

Turbocharging Fluid Antenna Multiple Access

Noor Waqar, Kai-Kit Wong, *Fellow, IEEE*, Chan-Byoung Chae, *Fellow, IEEE*, and Ross Murch, *Fellow, IEEE*

Abstract—Based on our current understanding, extreme massive access over the same physical channel is only possible if an extra-large multiple-input multiple-output (XL-MIMO) antenna is used at the base station (BS) and instantaneous channel state information (CSI) is known at the BS side for precoding design. This casts doubt on scalability and challenges in device-to-device situations in which there is not a centralized, optimized BS for transmitting the user signals. To address this problem, we revisit the massive connectivity challenge by considering the case where no CSI is available at the BS and no precoding is used. In this situation, inter-user interference (IUI) mitigation can only be performed at the user terminal (UT) side. Leveraging the position flexibility of fluid antenna system (FAS), we adopt a fluid antenna multiple access (FAMA) approach that exploits the interference signal fluctuation in the spatial domain. Specifically, we assume that we have N spatially correlated received signals per symbol duration from FAS. Our main approach uses a simple heuristic port shortlisting method that identifies promising ports to obtain favourable received signals that can be combined via maximum ratio combining (MRC) to form the received output signal for final detection. On top of this, a pre-trained deep joint source-channel coding (JSCC) scheme is employed, which together with a diffusion-based denoising model (MixDDPM) at the UT side, can improve the IUI immunity. We refer to the proposed scheme as *turbo FAMA*. Simulation results show that with a physical FAS size of 20 wavelengths at each UT transmitting quaternary phase shift keying (QPSK) symbols, fast FAMA can support 50 users while turbo FAMA can handle up to 200 users if the required symbol error rate (SER) is 10^{-2} . If a higher error tolerance is acceptable, say SER at 0.1, turbo FAMA can even serve up to 1000 users but fast FAMA is only able to handle 160 users, all remarkably achieved without CSI at the BS.

Index Terms—Fluid antenna systems, fast fluid antenna multiple access, denoising diffusion models, massive connectivity.

I. INTRODUCTION

A. Background

SPECTRUM is precious and the success of wireless communications technologies depends highly upon how much we can reuse our frequency bands [1]. In traditional orthogonal multiple access (OMA) approaches such as frequency-division multiple access (FDMA), there is an obvious trade-off between the number of supportable users and the bandwidth each user is allowed to have. Frequency reuse or more generally multiple

access on the same physical channel, on the other hand, hints at an exciting possibility of having both at the same time. Indeed, the highly anticipated sixth generation (6G) has the mission to deliver data rates of up to 100 Gbps, with peak rates reaching 1 Tbps, a connection density of 10^7 per km^2 , latency under 1 ms, and many other ambitious targets [2], [3], [4]. It is worth pointing out that latency in wireless networks is primarily due to the complexity of the protocol stack, and not the propagation delay. If multiple access techniques are highly complex, then latency will likely be high. Currently, there is a multiple access techniques (MAT) Industry Specification Group (ISG) set up in the European Telecommunications Standards Institute (ETSI) to explore novel multiuser access strategies [5].

In today's technologies, multiuser multiple-input multiple-output (MIMO) has been the most celebrated multiple access technique. The first multiuser MIMO work dates back in 2000 by Wong *et al.* [6] and since then, numerous results appeared, e.g., [7], [8], [9], [10], [11]. Multiuser MIMO permits multiple users to share the same physical channel by focusing the signal beams onto the desired users and placing nulls at unintended receivers—a technique called *precoding*. With decent channel state information (CSI) at the base station (BS) side, multiuser MIMO promises to handle as many users in the downlink as the number of BS antennas. In recent generations, multiuser MIMO has evolved into massive MIMO, promoting the use of an excessive number of BS antennas (typically 6 times the number of users) to simplify precoding design [12], [13]. The idea is elegant but in practice, massive MIMO reverts back to multiuser minimum mean-square-error (MMSE) precoding in the fifth generation (5G) [14] and the extra antennas are there to provide more degree-of-freedom (DoF) to avoid inter-user interference (IUI) more effectively. To realize 6G, a natural step will be to deploy even more BS antennas in the form of extra-large MIMO (XL-MIMO) [15] or a cell-free setup [16]. However, the overhead for CSI acquisition and the huge power consumption of a massive array call for a different design.

Besides MIMO, non-orthogonal multiple access (NOMA) and rate-splitting multiple access (RSMA) both have spurred enormous interest in recent years [17]. In NOMA, user signals are superimposed and IUI is subtracted at the user side using successive interference cancellation (SIC), see e.g., [18], [19], [20], [21]. NOMA does not scale nicely with the number of users since SIC becomes prohibitively complex very quickly. By contrast, RSMA manages the IUI by splitting the messages into common and private parts, and each user first decodes the common message by treating other signals as noise and then decodes its private message after subtracting the interference from the common message [22], [23]. Although SIC is still needed in RSMA, each user only needs to cancel the interference from the common message. As such, RSMA is believed to be more scalable than NOMA. Nevertheless, both schemes require CSI at the BS to perform power control as well as user clustering for the best performance. Due to high complexity,

The work of K. K. Wong is supported by the Engineering and Physical Sciences Research Council (EPSRC) under grant EP/W026813/1.

The work of R. Murch was supported by the Hong Kong Research Grants Council Area of Excellence Grant AoE/E-601/22-R.

The work of C. B. Chae was supported in part by the Korea government under NRF/Institute for Information and Communication Technology Planning and Evaluation (IITP) Grants 2021-0-00486 and 2022R1A5A1027646.

N. Waqar and K. K. Wong are with the Department of Electronic and Electrical Engineering, University College London, London WC1E 7JE, United Kingdom. K. K. Wong is also affiliated with Yonsei Frontier Laboratory, Yonsei University, Seoul, 03722, Korea.

C. B. Chae is with School of Integrated Technology, Yonsei University, Seoul, 03722, Korea.

R. Murch is with the Department of Electronic and Computer Engineering and Institute for Advanced Study (IAS), Hong Kong University of Science and Technology, Clear Water Bay, Hong Kong SAR, China.

Corresponding authors: Kai-Kit Wong and Chan-Byoung Chae.

the majority of results are limited to only a few users.

“Can space-division multiple-access be much simpler, at least conceptually?”

This was the question asked in [24] when pursuing a new multiple access method that is hopefully more straightforward, a lot more scalable, and does not require complex interference management, unlike multiuser MIMO, NOMA and RSMA. In this paper, we aim to shed light on the answer to a similar, but more specific question as follows:

“Can extreme massive multiple access be accomplished without CSI at the transmitter side, without centralized interference management and without SIC at the user side?”

Based on the results in [24], the answer is quite positive. Specifically, [24] proposed the concept of fluid antenna multiple access (FAMA). Through antenna position reconfiguration, a user can access the position, a.k.a. port, where the ratio of the desired signal energy to the energy of the sum-interference plus noise signal is maximized on a per-symbol basis, dealing with IUI entirely from the receiver side without interference management. This idea was later called fast FAMA [25].

The origin of the idea came from an earlier concept, referred to as fluid antenna system (FAS) [26], [27], [28]. FAS was first envisaged by Wong *et al.* in 2020 [29], [30] and FAS broadly represents the new generation of reconfigurable antennas that uses shape and position flexibility in antennas to empower the physical layer. Most recently, [31] and [32] explained the concept of FAS from the viewpoint of electromagnetics. A lot of opportunities for using artificial intelligence (AI) techniques in FAS were discussed in [33]. It is worth emphasizing that FAS is not restricted to any specific implementation methods. It can be realized by a variety of ways, such as liquid-based antennas [34], [35], [36], movable array [37], metamaterial-based antennas [38], and reconfigurable pixels [32], [39].

Following the seminal work in [29], [30], FAS has been extensively investigated in point-to-point communication scenarios. For instance, [40] proposed an eigenvalue-based channel model to facilitate diversity analysis for FAS channels. Then the authors in [41] presented a method to compute the diversity order based on the model in [40]. A spatial block-correlation model that maintains analytical tractability and approximates accurately the model for FAS channels was proposed in [42]. Subsequent research further extended the analysis to Nakagami fading channels [43], [44], α - μ fading channels [45] and even arbitrarily distributed fading channels [46]. Copula theory has also been shown to be a powerful tool for the analysis of FAS channels [47]. Recent efforts have also found FAS applied in MIMO systems and [48] uncovered the diversity-multiplexing trade-off of such systems. Finding the best antenna positions in FAS evidently requires CSI to be known. Efforts for estimating the CSI of FAS can be found in [49], [50], [51], [52].

A wide range of applications have already been illustrated to benefit from the deployment of FAS. For instance, FAS can introduce signal randomness to enhance physical layer security [53], [54], [55]. In [56], FAS was used to synergize full-duplex

communications. Reconfigurable intelligent surface (RIS) can also combine with FAS for promising performance [57], [58]. Moreover, in [59], [60], it was shown that FAS allows a new form of index modulation to trade-off diversity performance and communication rate. Recent trends also see efforts to use FAS for expanding the communication-sensing achievability region for integrated sensing and communication (ISAC) [61], [62], [63], [64]. It is worth stating that FAS broadly includes the emerging topics of movable antennas [65], flexible-position MIMO [66], and pinching antennas [67] as well.

B. Literature Review of FAMA

Going back to multiple access, FAMA indeed offers a very different understanding of how IUI can be mitigated—the one in which centralized interference management is not necessary and multipath fading becomes a desirable factor. In spite of a very positive outlook in [24], the authors did little to provide any clarity of how the best port can be found to maximize the ratio of the desired signal energy to the energy of the sum-interference plus noise signal on a per-symbol basis. Later in [25], an attempt was made to overcome this problem, but not without considerable performance loss. Furthermore, antenna positioning switching as fast as symbol rate, albeit doable via designs in [38], [32], is by no means trivial, which has led to the slow FAMA technique in [68], [69]. In slow FAMA, the FAS at each user only needs to switch its position once during each channel coherence period and remains unchanged unless the CSI changes. Finding the best position is also much easier as only the knowledge of the *average* signal-to-interference plus noise ratio (SINR) at the ports is required. There are also further efforts to exploit deep learning to assist the process of estimating the port SINR for slow FAMA [70], [71].

While slow FAMA is appealing for practical reasons, fast FAMA performs much better in terms of the number of users that can be supported [33]. For example, it is possible for fast FAMA to handle hundreds of users while slow FAMA is only able to deal with several users on the same physical channel.¹ To improve multi-access capability while maintaining simplicity, a variant of slow FAMA, coined as compact ultra massive array (CUMA), was proposed in [72]. This method shares the simplicity of slow FAMA to adapt the antenna position to only once each coherence time, but activates many ‘coherent’ ports for analogue signal mixing. CUMA can accommodate tens of co-channel users and works well even for channels with finite scattering if randomized RISs are deployed [73]. In addition, CUMA at the users combines well with multiuser MIMO at the BS to reduce the burden of CSI at the BS side for massive access [28, Section V-E]. On the other hand, recent attempts have also found channel coding to be a crucial ingredient to strengthen slow FAMA [74], [75]. Opportunistic scheduling is another technique to leverage slow FAMA for enhancing the network performance [76]. Interestingly, although FAMA can be a standalone multiple access scheme, FAS can be used to facilitate more reliable operations of NOMA [77] and RSMA [78] for realizing their capacity advantages.

¹Note that this comment assumes a not so ambitious SINR target for each user, and the exact number of supportable users reduces greatly if the required SINR at the users increases.

C. Challenges and Contributions

In summary, CUMA and slow FAMA are practically attractive but they are not designed to handle the IUI for supporting hundreds of users or more on the same channel. This is what fast FAMA promises to do theoretically but it is not known to be practically achievable. Yet, massive access methods without the need for CSI at the BS side nor centralized optimization are profoundly important for 6G and beyond systems. In this paper, our ambition is to devise a FAMA scheme that performs better than fast FAMA and is practically realizable.

To achieve this goal, it is natural to resort to AI techniques due to the complexity and unknown nature of the problem. In fact, there is an increasing trend of employing AI to overcome complex wireless optimization problems in real time, and they are poised to play a major role in 6G [79]. AI-driven methods to learn intricate interference patterns, make rapid data-driven decisions for resource allocation, beamforming, and decoding, overcoming the limitations of traditional model-based methods have previously been proposed [80], [81], [82]. For example, deep neural networks (DNNs) have been successfully applied to multiuser detection, demonstrating robustness to interference and channel distortions, while reinforcement learning has been used for dynamic spectrum access and power control in interference-limited networks [83], [84], [85].

In the context of FAMA, AI methods are capable of learning to map observed interference metrics to optimal antenna port selections or predict future interference levels for proactive port switching. In this regard, generative AI models, such as generative adversarial networks and diffusion models, which iteratively denoise random inputs to recover clean signals, are particularly promising [86], [87], [88]. Recently, these models have been employed for channel estimation, signal detection, and constellation shaping [89], [90], [91], [92]. Their ability to handle high-dimensional data and operate in real time after training motivates our work to integrate an AI-driven diffusion module into FAMA for enhanced interference suppression.

Another promising approach to improve resilience of wireless links is joint source-channel coding (JSCC) [93], [94]. JSCC jointly designs the transmitted signals by taking into account both the source characteristics (for compression) and the channel conditions (for error correction), thereby introducing redundancy that mirrors the source content and significantly enhances system robustness [95], [96]. In interference-prone environments, JSCC can exploit the inherent structure of both the source and the channel interference to achieve robustness unattainable by classical error-correcting codes designed for independent and identically distributed (i.i.d.) errors. For this reason, in this work, we will leverage JSCC by integrating a diffusion-based denoiser into the decoder, thereby treating the interference-plus-noise signal as a contaminant to be removed during joint detection and decoding. Instead of encoding each symbol independently, our system embeds redundancy across symbols, so that even if some symbols are heavily corrupted by interference as well as noise, the overall message can still be reliably recovered. This integration of JSCC with FAMA not only harnesses the spatial diversity of fluid antennas but further provides robust diversity, resulting in an interference-

resilient communication scheme. Motivated by the limitations of current FAMA schemes and recent advances in denoising diffusion and JSCC models, we leverage AI-driven diffusion-based denoising and deep JSCC to improve FAMA.

Specifically, in this paper, we develop a new FAMA scheme that heuristically shortlists a few promising ports that produce an output signal using maximum ratio combining (MRC) in the computational domain. This scheme is integrated with an AI-driven interference mitigation framework and JSCC. The proposed scheme is referred to as turbo FAMA. Our proposed turbo FAMA framework uniquely integrates fluid antenna's spatial diversity with deep JSCC and diffusion-based denoising, comprehensively addressing the massive access challenge in CSI-free downlink channels. In particular, fluid antennas exploit interference fluctuations across spatial antenna positions, significantly reducing dominant interferers via port shortlisting and coherent MRC. However, residual structured interference persists, calling for advanced coding strategies. In our case, deep JSCC jointly optimizes source compression and interference-aware error correction, embedding redundancy specifically tailored to structured interference. Also, the MixDDPM diffusion-based denoising model further eliminates the remaining non-Gaussian, heavy-tailed interference patterns unaddressed by traditional linear methods. Overall, by addressing interference at multiple complementary stages, physical-layer diversity, learned denoising, and JSCC, our framework "turbocharges" conventional fluid antenna technology, achieving levels of interference mitigation and reliability that neither signal processing nor JSCC alone could provide.

The contributions are summarized as follows.

- First, we propose a new FAMA architecture that shortlists several promising ports for MRC.
- Moreover, we develop an AI-driven interference mitigation framework based on a mix diffusion-based denoising model, which adaptively learns and suppresses residual interference which exhibits uncertain characteristics.
- Also, we integrate a deep JSCC scheme into the FAMA framework, ensuring robust end-to-end (e2e) data transmission by jointly optimizing source and channel coding in interference-prone environments.
- Simulation results in terms of symbol error rate (SER) will demonstrate that the proposed turbo FAMA scheme outperforms greatly fast FAMA with ideal port selection and can support an extremely large number of users on the same physical channel without CSI at the BS.

The rest of this paper is organized as follows. Section II introduces the system model, describing the channel model of FAS and the interference network model. Section III presents our port-selection strategies while Section IV details the proposed diffusion-based denoising and its integration with the deep JSCC scheme. In Section V, we provide the simulation results. Finally, Section VI concludes the paper.

II. SYSTEM MODEL

Consider a single-cell downlink system in which a BS with N_t spatially distributed antennas serves U single-antenna user

terminals (UTs), with $N_t = U$ ². Each BS antenna is dedicated to serving one UT, such that the u -th BS antenna primarily transmits data intended for UT u , for $u = 1, 2, \dots, U$. At the user side, each UT adopts an FAS comprising K discrete ports arranged linearly along an aperture of length $W\lambda$, where λ denotes the carrier wavelength.³ The FAS enables fast position switching of the active port among the K positions so that it is reasonable to assume that with FAS, UT u can access the received signals from all K positions in the given space.⁴

A. Signal Model

Let $s_u \in \mathbb{C}$ be the transmitted symbol from the u -th BS antenna, satisfying $\mathbb{E}[|s_u|^2] = \sigma_s^2$ and assumed independent across u . The complex channel gain from the $\bar{u}$ -th BS antenna to the k -th port of user u is denoted by $g_k^{(\bar{u},u)}$. The received baseband signal at the k -th port of UT u is expressed as

$$r_k^{(u)} = g_k^{(u,u)} s_u + \sum_{\substack{\bar{u}=1 \\ \bar{u} \neq u}}^U g_k^{(\bar{u},u)} s_{\bar{u}} + \eta_k^{(u)}, \quad (1)$$

where $\eta_k^{(u)} \sim \mathcal{CN}(0, \sigma_\eta^2)$ is the circularly symmetric complex Gaussian (CSCG) noise at port k of UT u , independent across u and k . The average received signal-to-noise ratio (SNR) per port for the desired signal is defined as

$$\Gamma_u \triangleq \frac{\Omega_u \sigma_s^2}{\sigma_\eta^2}, \quad (2)$$

where $\Omega_u = \mathbb{E}[|g_k^{(u,u)}|^2]$ is the average channel gain of the desired signal, assumed constant across ports k for UT u .

B. Channel Model

The value of K is usually large and hence the ports in FAS are closely spaced, resulting in high spatial correlation among the channel coefficients within the length of $W\lambda$. For UT u , define the channel vector from the u -th BS antenna to the K ports as $\mathbf{g}_{u,u} = [g_1^{(u,u)}, g_2^{(u,u)}, \dots, g_K^{(u,u)}]^T$. We assume that $\mathbf{g}_{u,u} \sim \mathcal{CN}(\mathbf{0}, \Omega_{u,u} \mathbf{J})$, where $\Omega_{u,u} = \mathbb{E}[|g_k^{(u,u)}|^2]$ is the average channel gain from BS antenna u to UT u 's ports, and $\mathbf{J} \in \mathbb{C}^{K \times K}$ is the spatial correlation matrix. The channel vectors $\mathbf{g}_{u,u}$ are assumed independent across different u . Also, the (n, m) -th element of $\mathbf{J}$ is given by

$$[\mathbf{J}]_{n,m} = J_0 \left(2\pi \frac{|n-m|}{K-1} W \right), \quad (3)$$

²The condition of $N_t = U$ is considered naturally because in our work, no precoding is adopted and each BS antenna is dedicated to transmitting one UT's signal. This also represents a typical fully-loaded downlink scenario. However, this condition is not fundamentally required, and our proposed framework can readily adapt to scenarios where $N_t \neq U$. For instance, if $N_t < U$, practical implementation would necessitate user scheduling, grouping, or superimposed transmission schemes, potentially increasing interference intensity. Conversely, if $N_t > U$, excess antennas could be utilized through diversity coding, combining or spatial grouping, although exploiting such strategies without CSI-based precoding is beyond the scope of this paper.

³This work can be easily extended to the case if each UT is equipped with a two-dimensional (2D) FAS.

⁴This might mean that each UT switches its port K times in each symbol duration, or $\frac{K}{n_{\text{RF}}}$ times if it has n_{RF} radio-frequency (RF) chains. Evidently, it is also possible to reduce the number of switching times per symbol duration by exploiting the spatial correlation between the positions using AI techniques, which was studied in [97]. However, we defer a full investigation of how this technique works in the context of FAMA to future work.

in which $J_0(\cdot)$ is the zeroth-order Bessel function of the first kind, modeling correlation as a function of port separation.

To generate $\mathbf{g}_{u,u}$, we employ the eigenvalue decomposition of $\mathbf{J} = \mathbf{Q} \mathbf{\Lambda} \mathbf{Q}^H$, where $\mathbf{Q}$ is a unitary matrix of eigenvectors, and $\mathbf{\Lambda} = \text{diag}(\lambda_1, \lambda_2, \dots, \lambda_K)$ contains the eigenvalues satisfying $\lambda_1 \geq \lambda_2 \geq \dots \geq \lambda_K \geq 0$. Thus,

$$\mathbf{g}_{u,u} = \sqrt{\Omega_{u,u}} \mathbf{Q} \mathbf{\Lambda}^{\frac{1}{2}} \mathbf{w}_{u,u}, \quad (4)$$

in which $\mathbf{w}_{u,u} \sim \mathcal{CN}(\mathbf{0}, \mathbf{I}_K)$ is a vector of independent CSCG entries. This ensures $\mathbb{E}[\mathbf{g}_{u,u} \mathbf{g}_{u,u}^H] = \Omega_{u,u} \mathbf{J}$, preserving the desired correlation structure.

C. Symbol Level Port Switching

Fast FAMA leverages symbol-level port selection to maximize the instantaneous SINR. In particular, fast FAMA adapts the active port k_u^* for each symbol based on real-time channel conditions in the data-dependent interference-plus-noise signal. Define the instantaneous SINR at port k of UT u as

$$\gamma_k^{(u)} = \frac{|g_k^{(u,u)}|^2 |s_u|^2}{\left| \sum_{\bar{u} \neq u} g_k^{(\bar{u},u)} s_{\bar{u}} + \eta_k^{(u)} \right|^2}. \quad (5)$$

Ideally, UT u should select

$$k_u^* = \arg \max_{1 \leq k \leq K} \gamma_k^{(u)}. \quad (6)$$

Since $|s_u|^2$ is constant across ports for a given symbol, this is equivalent to

$$k_u^* = \arg \max_{1 \leq k \leq K} \frac{|g_k^{(u,u)}|^2}{\left| \sum_{\bar{u} \neq u} g_k^{(\bar{u},u)} s_{\bar{u}} + \eta_k^{(u)} \right|^2}. \quad (7)$$

The efficacy of fast FAMA arises from exploiting deep fades in the data-dependent interference term $\sum_{\bar{u} \neq u} g_k^{(\bar{u},u)} s_{\bar{u}} + \eta_k^{(u)}$. For large U , central limit theorem (CLT) applies and this term behaves like a complex Gaussian random variable. As a result, the magnitude $\left| \sum_{\bar{u} \neq u} g_k^{(\bar{u},u)} s_{\bar{u}} + \eta_k^{(u)} \right|$ is Rayleigh distributed, which is well known to suffer from deep fades with significant likelihood. These deep fades enable the fast FAMA mechanism to choose a port k in which the aggregate interference signal is momentarily suppressed, boosting the instantaneous SINR.

Despite its effectiveness to neutralize interference by natural fading, estimation of the instantaneous, data-dependent SINR for port selection is non-trivial. This issue was tackled in [25] but considerable performance degradation was observed.

III. HEURISTIC-BASED PORT SHORTLISTING

Accurately estimating the instantaneous SINR at the ports to determine the best port is very challenging. Here, we present several heuristics that identify ports with desirable properties for processing. The heuristics are simple and realizable.

A. Port Shortlisting

Given that each user, say UT u , possesses the knowledge of the desired channel $g_k^{(u,u)}$ [49], [50], [51], [52], the predicted received power of the desired signal can be expressed as

$$P_{\text{desired}}^{(u)}(k) = |g_k^{(u,u)}|^2 \sigma_s^2. \quad (8)$$

A straightforward suboptimal selection rule involves identifying those ports k whose power $|r_k^{(u)}|^2$ most closely matches (8). Large deviations from this value suggest either a strong interference component or destructive fading of the desired signal. As such, we define the deviation metric as

$$\Delta_k^{(u)} \triangleq \left| |r_k^{(u)}|^2 - P_{\text{desired}}^{(u)}(k) \right|. \quad (9)$$

Ports are then ranked in ascending order of $\Delta_k^{(u)}$.

To quantify the reliability of (9), consider the received signal $r_k^{(u)} = g_k^{(u,u)} s_u + I_k^{(u)} + \eta_k^{(u)}$, where $I_k^{(u)} \triangleq \sum_{\bar{u} \neq u} g_k^{(\bar{u},u)} s_{\bar{u}}$. For large U , $I_k^{(u)}$ approximates a complex Gaussian random variable by CLT, with zero mean and approximate variance of $\sigma_I^2 = (U-1)\sigma_s^2\Omega_{\text{cross}}$, where $\Omega_{\text{cross}} = \mathbb{E}[|g_k^{(\bar{u},u)}|^2]$. The received power $|r_k^{(u)}|^2$ follows a noncentral chi-squared distribution with 2 DoF, parameterized by the desired signal power $|g_k^{(u,u)}|^2\sigma_s^2$ and the interference-plus-noise variance $\sigma_I^2 + \sigma_\eta^2$. Thus, the probability that $|r_k^{(u)}|^2$ deviates from $P_{\text{desired}}^{(u)}(k)$ by less than δ can be bounded by using

$$\begin{aligned} P(\Delta_k^{(u)} < \delta) \\ = P\left(P_{\text{desired}}^{(u)}(k) - \delta < |r_k^{(u)}|^2 < P_{\text{desired}}^{(u)}(k) + \delta\right). \end{aligned} \quad (10)$$

Given that $|r_k^{(u)}|^2 \sim (\sigma_I^2 + \sigma_\eta^2) \chi_2^2\left(\frac{P_{\text{desired}}^{(u)}(k)}{\sigma_I^2 + \sigma_\eta^2}\right)$, the probability can thus be represented in terms of the non-central chi-squared cumulative density function (CDF) as [103], [104]

$$\begin{aligned} P(\Delta_k^{(u)} < \delta) = \mathcal{F}_{\chi_2^2\left(\frac{P_{\text{desired}}^{(u)}(k)}{\sigma_I^2 + \sigma_\eta^2}\right)}\left(\frac{P_{\text{desired}}^{(u)}(k) + \delta}{\sigma_I^2 + \sigma_\eta^2}\right) \\ - \mathcal{F}_{\chi_2^2\left(\frac{P_{\text{desired}}^{(u)}(k)}{\sigma_I^2 + \sigma_\eta^2}\right)}\left(\frac{P_{\text{desired}}^{(u)}(k) - \delta}{\sigma_I^2 + \sigma_\eta^2}\right). \end{aligned} \quad (11)$$

This bound provides some theoretical insight for selecting ports with minimal deviation.

To refine this approach, we normalize the deviation by the predicted desired-signal power, yielding

$$\Delta_k^{(u)} = \frac{\left| |r_k^{(u)}|^2 - |g_k^{(u,u)}|^2\sigma_s^2 \right|}{|g_k^{(u,u)}|^2\sigma_s^2}, \quad (12)$$

where the relative difference governs the decision. This normalization prevents ports with strong desired channels from being unfairly penalized due to larger absolute deviations.

Additionally, in systems with numerous ports, pre-filtering weak desired-channel gains improves efficiency. To do so, we define the maximum desired-channel gain as

$$|g_{\text{max}}^{(u,u)}|^2 \triangleq \max_{1 \leq k \leq K} |g_k^{(u,u)}|^2, \quad (13)$$

and select a candidate set

$$\mathcal{K}^{(u)} = \left\{ k : |g_k^{(u,u)}|^2 \geq \gamma_{\text{th}} |g_{\text{max}}^{(u,u)}|^2 \right\}, \quad (14)$$

where $0 < \gamma_{\text{th}} < 1$ is a threshold parameter. This set excludes ports with negligible desired signals, reducing the risk of selecting interference-dominated ports.

The third rule is regarding spatial diversity maximization (SDM). From $\mathcal{K}^{(u)}$, we select K_{sel} ports to maximize spatial diversity within the aperture $W\lambda$. The optimal set $\mathcal{K}_{\text{SDM}}^{(u)}$ solves the following combinatorial optimization:

$$\mathcal{K}_{\text{SDM}}^{(u)} = \arg \max_{\substack{\mathcal{K} \subset \mathcal{K}^{(u)} \\ |\mathcal{K}| = K_{\text{sel}}}} \min_{k, \ell \in \mathcal{K}} \delta(k, \ell), \quad (15)$$

where the normalized port separation is given by

$$\delta(k, \ell) = \frac{|k - \ell|}{K - 1} W. \quad (16)$$

This max-min formulation ensures that the selected ports are as far apart as possible, decorrelating their channels for more diversity benefits. This selection process is illustrated in Fig. 1 in which SDM is not applied but we impose a fixed minimum spacing between adjacent selected ports $d \times K$ with $d = 0.05$.

The proposed port shortlisting algorithm exhibits computational complexity of approximately $O(K \log K)$, with linear complexity $O(K)$ for deviation metric computation and $O(K \log K)$ for sorting operations, which remains computationally manageable even for large K . Overall, the heuristic-based port shortlisting scheme considering all the procedures above is given in Algorithm 1.

Algorithm 1 Port Shortlisting Process

Require: Knowledge of CSI $\{g_k^{(u,u)}\}$, threshold γ_{th}

Ensure: Selected port set $\mathcal{K}_{\text{SDM}}^{(u)}$

- 1: Compute predicted desired power for all ports using (8)
 - 2: Compute deviation metric $\Delta_k^{(u)}$ for all ports using (12)
 - 3: Rank ports in ascending order of $\Delta_k^{(u)}$
 - 4: Apply desired power filtering to the ranked list using (14)
 - 5: Solve the SDM from the filtered list using (15)
-

B. MRC for Shortlisted Ports

After selecting K_{sel} ports using Algorithm 1, MRC is used to coherently combine the signals for maximizing the effective SNR. Specifically, MRC weights the signal of each shortlisted port by the conjugate of its normalized channel gain to produce the output signal for UT u as

$$y_u = \sum_{k \in \mathcal{K}_{\text{SDM}}^{(u)}} \frac{\left(g_k^{(u,u)}\right)^*}{\sqrt{\Omega_u}} r_k^{(u)}. \quad (17)$$

MRC introduces minimal complexity, without requiring any knowledge of the interference.⁵ It is worth stating that MRC is performed in the computational domain rather than in the signal domain, meaning that using FAS, fast port switching per symbol period allows for the received signals to be obtained on one or a small number of RF chains. Once all the received signals $\{r_k^{(u)}\}$ are obtained, they are recorded and processed in the computational domain to produce y_n for decoding.

⁵While linear interference-suppression methods such as zero-forcing (ZF) or minimum mean square error (MMSE) combining theoretically offer improved interference mitigation, they are inapplicable here. The reason is that their implementation requires extensive CSI from the interfering users, significantly increasing complexity and training overhead. In massive-access scenarios, acquiring such interference CSI is practically infeasible.

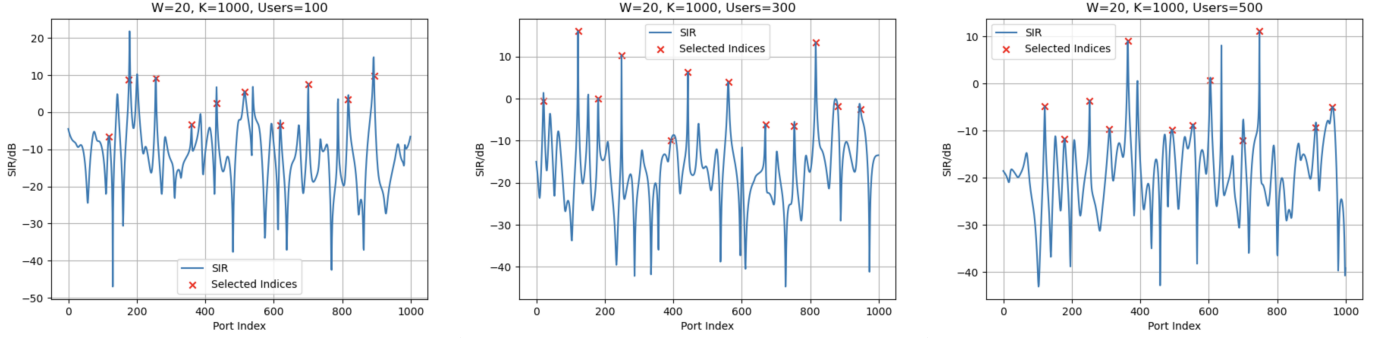


Fig. 1. Illustration of the port shortlisting process with $\gamma_{th} = 0.6$, $K = 1000$ ports and $W = 20$, for 100, 300, and 500 users. The blue lines represent the instantaneous SINR across all port indices, while the red 'x' markers indicate the selected ports that satisfy the shortlisting criteria (without SDM).

IV. DIFFUSION-BASED DENOISING FOR RESIDUAL INTERFERENCE REMOVAL

The proposed heuristics mitigate dominant interference by shortlisting ports with favorable conditions based on the set criteria. However, the residual interference plus noise at the combined signal output after MRC, denoted as η_{res} , remains. This residual interference exhibits non-Gaussian characteristics such as heavier tails and multi-modal distributions due to partial constructive alignment or destructive misalignment of multiuser signals. Classical detectors, such as linear or MMSE estimators, rely on Gaussian noise assumptions and thus incur significant performance degradation under these conditions. To overcome this, here, we propose a diffusion-based denoising framework that leverages generative modeling to learn and remove structured residual interference and noise, enhancing symbol detection accuracy in the FAMA system.

The motivation behind adopting diffusion-based denoising stems from the dynamic and highly sophisticated nature of this data-dependent residual interference, η_{res} . Symbol-level port switching introduces chaotic, but correlated variations in interference patterns, rendering traditional model-based approaches requiring precise CSI or multiuser inversion infeasible. On the contrary, diffusion models, operating in a data-driven manner, adaptively capture the statistical distribution of clean signals and interference without explicit CSI, thereby accommodating non-Gaussianity and multi-modality.

A. Denoising Diffusion Probabilistic Models (DDPM)

DDPM [98] employs a two-phase process, a *forward diffusion* phase that corrupts a clean signal with noise over T steps, and a *reverse denoising* phase that reconstructs the original signal from its noisy version. It should be noted that T steps here may mean the number of iterations required for training the model or that for iterative processing when applied.

1) *Forward diffusion process*: Let $\mathbf{x}_0 \in \mathbb{C}^n$ represent the clean received symbol vector after MRC in which n denotes the number of signal samples in time (i.e., $y_u(1), \dots, y_u(n)$). For mathematical convenience, we stack its real and imaginary parts, treating $\mathbf{x}_0 \in \mathbb{R}^{2n}$. Define a noise schedule $\{\beta_t\}_{t=1}^T$, where $\beta_t \in (0, 1)$, $\alpha_t = 1 - \beta_t$, and $\bar{\alpha}_t = \prod_{i=1}^t \alpha_i$. The forward diffusion process is performed by

$$q(\mathbf{x}_t | \mathbf{x}_{t-1}) = \mathcal{N}(\mathbf{x}_t; \sqrt{\alpha_t} \mathbf{x}_{t-1}, \beta_t \mathbf{I}), \text{ for } 1 \leq t \leq T, \quad (18)$$

with the joint distribution factorized as

$$q(\mathbf{x}_{1:T} | \mathbf{x}_0) = \prod_{t=1}^T q(\mathbf{x}_t | \mathbf{x}_{t-1}), \quad (19)$$

where $\mathbf{x}_{1:T}$ is the shorthand for $\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_T$. This process incrementally adds Gaussian noise, scaling the signal by $\sqrt{\alpha_t}$ and introducing variance β_t , to allow direct sampling

$$\mathbf{x}_t = \sqrt{\alpha_t} \mathbf{x}_0 + \sqrt{1 - \alpha_t} \boldsymbol{\epsilon}, \text{ where } \boldsymbol{\epsilon} \sim \mathcal{N}(\mathbf{0}, \mathbf{I}). \quad (20)$$

After T steps, we have $\mathbf{x}_T \approx \mathcal{N}(\mathbf{0}, \mathbf{I})$, effectively erasing the original signal's structure and containing only noise.

2) *Reverse denoising process*: The reverse process, parameterized by p_θ , approximates the posterior

$$p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t) = \mathcal{N}(\mathbf{x}_{t-1}; \boldsymbol{\mu}_\theta(\mathbf{x}_t), \sigma_t^2 \mathbf{I}), \quad (21)$$

where $\boldsymbol{\mu}_\theta(\mathbf{x}_t)$ is predicted by a neural network, and σ_t^2 is either fixed or learned. Typically, we have

$$\boldsymbol{\mu}_\theta(\mathbf{x}_t) = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \frac{\beta_t}{\sqrt{1 - \alpha_t}} \boldsymbol{\epsilon}_\theta(\mathbf{x}_t) \right), \quad (22)$$

with $\boldsymbol{\epsilon}_\theta(\mathbf{x}_t)$ estimating the noise $\boldsymbol{\epsilon}$.

3) *Training objective*: Training minimizes a simplified variational bound, reducing to noise prediction

$$\mathcal{L}_{\text{DDPM}} = \mathbb{E}_{t, \mathbf{x}_0, \boldsymbol{\epsilon}} \left[\|\boldsymbol{\epsilon} - \boldsymbol{\epsilon}_\theta(\mathbf{x}_t)\|^2 \right]. \quad (23)$$

Sampling involves

$$\begin{cases} \mathbf{x}_T \sim \mathcal{N}(\mathbf{0}, \mathbf{I}), \\ \mathbf{x}_{t-1} \sim p_\theta(\mathbf{x}_{t-1} | \mathbf{x}_t), \text{ for } t = T, \dots, 1. \end{cases} \quad (24)$$

In symbol-level FAMA (e.g., fast FAMA), $\mathbf{x}_0$ represents the desired signal, and η_{res} is treated as the initial corruption. But DDPM Gaussian noise assumption limits its ability to model multi-modal interference, prompting an extension.

B. Mixture-based Denoising Diffusion (MixDDPM)

To address the limitations of standard DDPMs in modeling the complex, multi-modal interference distributions encountered in symbol-level FAMA systems, we introduce the MixDDPM, as shown in Fig. 2 [99], [100]. This framework extends the traditional model by incorporating Gaussian mixture noise, enabling explicit representation of interference modes (e.g., constructive/destructive alignment) in fast FAMA systems. The

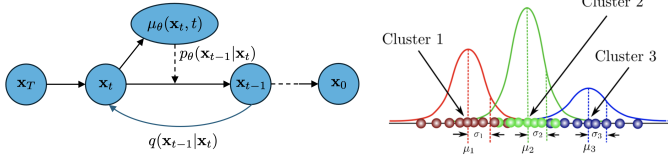


Fig. 2. The MixDDPM framework, depicting the forward diffusion process $q(\mathbf{x}_t|\mathbf{x}_{t-1})$ that gradually corrupts the clean signal and the reverse denoising process $p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t)$ that iteratively recovers $\mathbf{x}_0$.

key innovation lies in the ability to capture the structured nature of residual interference, which exhibits non-Gaussian characteristics due to the superposition of multiuser signals.

1) *Forward process with mixture noise*: The forward diffusion process in MixDDPM is defined as

$$q(\mathbf{x}_t|\mathbf{x}_{t-1}) = \sum_{j=1}^J \omega_j \mathcal{N}(\mathbf{x}_t; \sqrt{\alpha_t} \mathbf{x}_{t-1} + \sqrt{\beta_t} \mathbf{c}_j, \beta_t \mathbf{I}), \quad (25)$$

in which $\{\mathbf{c}_j\}_{j=1}^J$ are the mixture centers, ω_j are the mixture weights satisfying $\sum_{j=1}^J \omega_j = 1$, and J denotes the number of modes. All of these parameters can be fixed or learned. The marginal distribution at step t follows (20) as before.

The Gaussian-mixture diffusion process allows for closed-form sampling of $\mathbf{x}_t$ at any timestep t . To do so, let $\gamma_{t,i} = \beta_i \prod_{j=i+1}^t \alpha_j$ for $i \in [1, t-1]$ and $\gamma_{t,t} = \beta_t$. Considering the iterates of $\mathbf{x}_t$, we then have

$$\mathbf{x}_t = \sqrt{\alpha_t} \mathbf{x}_0 + \sqrt{\sum_{j=1}^J \gamma_{t,j} \sigma_j^2} \bar{\mathbf{z}}_t + \sum_{j=1}^J \sqrt{\gamma_{t,j}} \mathbf{c}_j, \quad (26)$$

where $\bar{\mathbf{z}}_t \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ and j denotes the Gaussian component selected at timestep t . In the sequel, when required, we will use i_t to replace j to highlight the time dependence. This closed-form expression simplifies the sampling process and enables efficient implementation. This is described in Algorithm 2.

Algorithm 2 MixDDPM Forward Diffusion Process

Require: Input signal $\mathbf{x}_0$, noise schedule $\{\beta_t\}$, mixture components $\{\mathbf{c}_j, \omega_j\}$, timesteps T

Ensure: Noisy signal $\mathbf{x}_T$

- 1: **for** $t = 1$ to T **do**
- 2: Sample index $j \sim \text{Categorical}(\omega_1, \dots, \omega_K)$
- 3: Sample noise $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$
- 4: Update signal by (26)
- 5: **end for**

2) *Reverse process with mixture components*: The reverse denoising process is parameterized by

$$p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t) = \sum_{j=1}^K p_\theta(j|\mathbf{x}_t) \mathcal{N}(\mathbf{x}_{t-1}; \boldsymbol{\mu}_\theta^{(j)}(\mathbf{x}_t), \boldsymbol{\Sigma}_\theta^{(j)}(\mathbf{x}_t)), \quad (27)$$

where $p_\theta(j|\mathbf{x}_t)$ classifies the interference mode, and $\boldsymbol{\mu}_\theta^{(j)}$ and $\boldsymbol{\Sigma}_\theta^{(j)}$ are learned parameters for each component. The mean prediction incorporates the mixture centers

$$\boldsymbol{\mu}_\theta^{(j)}(\mathbf{x}_t) = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \sqrt{\beta_t} \mathbf{c}_{i_t} - \frac{\beta_t}{\sqrt{1-\alpha_t}} \boldsymbol{\epsilon}_\theta^{(j)}(\mathbf{x}_t) \right) + \mathbf{c}_j, \quad (28)$$

where $\boldsymbol{\epsilon}_\theta^{(j)}(\mathbf{x}_t)$ estimates the noise specific to component j . This is summarized in Algorithm 3.

Algorithm 3 MixDDPM Reverse Denoising Process

Require: Noisy signal $\mathbf{x}_T$, trained model p_θ , timesteps T

Ensure: Denoised signal $\hat{\mathbf{x}}_0$

- 1: Initialize $\mathbf{x}_T$
- 2: **for** $t = T$ downto 1 **do**
- 3: Predict noise $\boldsymbol{\epsilon}_\theta^{(j)}(\mathbf{x}_t, t)$ for all components j
- 4: Compute mixture weights $p_\theta(j|\mathbf{x}_t)$
- 5: Sample component index $j \sim p_\theta(j|\mathbf{x}_t)$
- 6: Update signal by $\mathbf{x}_{t-1} = \boldsymbol{\mu}_\theta^{(j)}(\mathbf{x}_t)$ in (28)
- 7: **end for**

3) *Training objective and Kullback-Leibler (KL) divergence*: The loss function combines noise prediction and interference mode classification, given by

$$\mathcal{L}_{\text{MixDDPM}} = \mathbb{E}_{t, \mathbf{x}_0, \epsilon, j} \left[-\log p_\theta(j|\mathbf{x}_t) + \left\| \epsilon - \boldsymbol{\epsilon}_\theta^{(j)}(\mathbf{x}_t) \right\|^2 \right], \quad (29)$$

where j is sampled from the categorical distribution defined by ω_j . The training procedure minimizes the variational lower bound (VLB) of the negative log-likelihood [101]

$$\begin{aligned} L_{\text{VLB}} = & -\log p_\theta(\mathbf{x}_0|\mathbf{x}_1) \\ & + \sum_{t>1} \mathbb{E}_{i_{1:t-1}|\mathbf{x}_t, \mathbf{x}_0} [D_{\text{KL}}(q(i_t|\mathbf{x}_t, \mathbf{x}_0, i_{1:t-1}) \| p_\theta(i_t|\mathbf{x}_t))] \\ & + \sum_{t>1} \mathbb{E}_{i_{1:t}|\mathbf{x}_t, \mathbf{x}_0} [D_{\text{KL}}(q(\mathbf{x}_{t-1}|\mathbf{x}_t, \mathbf{x}_0, i_{1:t}) \| p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t, i_t))] \\ & + D_{\text{KL}}(q(\mathbf{x}_T|\mathbf{x}_0) \| p(\mathbf{x}_T)), \end{aligned} \quad (30)$$

in which $D_{\text{KL}}(P \| Q)$ represents the KL divergence. The KL divergence terms can be computed in closed form as

$$\begin{aligned} & \mathbb{E}_{i_{1:t}|\mathbf{x}_t, \mathbf{x}_0} [D_{\text{KL}}(q(\mathbf{x}_{t-1}|\mathbf{x}_t, \mathbf{x}_0, i_{1:t}) \| p_\theta(\mathbf{x}_{t-1}|\mathbf{x}_t, i_t))] \\ = & \mathbb{E}_{i_t|\mathbf{x}_t, \mathbf{x}_0} \left[\left\| \begin{aligned} & \frac{\sigma_t^2}{2\sigma_t^2(i_t)} + \frac{1}{2\sigma_t^2(i_t)} \times \\ & \mathbb{E}_{i_{1:t-1}|\mathbf{x}_t, \mathbf{x}_0, i_t} [\boldsymbol{\mu}_\theta(\mathbf{x}_t, \mathbf{x}_0, i_{1:t-1})] \\ & - \boldsymbol{\mu}_\theta(\mathbf{x}_t, t, i_t) \\ & + \log \frac{\sigma_t(i_t)}{\sigma_t} \end{aligned} \right\|_2^2 \right] \\ & + C, \end{aligned} \quad (31)$$

where C is a constant independent of θ .

To predict $\boldsymbol{\mu}_\theta(\mathbf{x}_t)$, we use the parameterization

$$\boldsymbol{\mu}_\theta(\mathbf{x}_t, i_t) = \frac{1}{\sqrt{\alpha_t}} \left(\mathbf{x}_t - \sqrt{\beta_t} \mathbf{c}_{i_t} - \frac{\beta_t}{\sqrt{1-\alpha_t}} \boldsymbol{\epsilon}_\theta(\mathbf{x}_t) \right), \quad (32)$$

where $\boldsymbol{\epsilon}_\theta(\mathbf{x}_t)$ denotes a neural network predicting the noise component. This parameterization ensures that the model can effectively learn the structure of interference modes while maintaining computational efficiency.

The Gaussian-mixture diffusion process permits simple, closed-form sampling of $\mathbf{x}_t$ at any timestep t . Furthermore, the use of fast ordinary differential equation (ODE) solvers during the reverse process accelerates convergence, making the framework suitable for real-time applications in symbol-level FAMA systems. The explicit modeling of interference modes

through mixture components enhances the model's ability to adapt to fast changing channel conditions.

In the proposed FAMA, MixDDPM addresses the challenge of residual interference by explicitly modeling the multi-modal nature of interference. The mixture components correspond to different interference alignment scenarios (e.g., constructive versus destructive), enabling the denoiser to adaptively remove structured interference. This integration improves symbol detection accuracy in interference-limited regimes, a key advantage over traditional Gaussian-based denoising methods.

C. JSCC with Denoising Diffusion

To enhance the robustness of the proposed FAMA system, we combine the diffusion-based denoising pipeline with a deep JSCC framework (encoder and decoder). The JSCC encoder transforms digital data symbols such as quaternary phase shift keying (QPSK) into a latent space representation optimized for interference resilience, while the decoder reconstructs the original symbols from the denoised latent vector. This end-to-end optimized system explicitly addresses residual multiuser interference and improves symbol estimation accuracy.

The JSCC encoder $\mathcal{E}_\phi(\cdot)$ transforms the transmitted QPSK symbol vector $\mathbf{s}_u \in \mathbb{C}^n$ for UT u into a latent representation $\mathbf{z}_u \in \mathbb{R}^{2n}$ (with n being the block size) so that

$$\mathbf{z}_u = \mathcal{E}_\phi(\mathbf{s}_u), \quad (33)$$

where ϕ represents the encoder's learnable parameters. The encoder employs convolutional neural layers to exploit local correlations in the symbol vector

$$\mathbf{h}^{(l)} = \sigma \left(\mathbf{W}^{(l)} \mathbf{h}^{(l-1)} + \mathbf{b}^{(l)} \right), \text{ for } l = 1, \dots, L, \quad (34)$$

with the input $\mathbf{h}^{(0)} = \text{Re}(\mathbf{s}_u) \oplus \text{Im}(\mathbf{s}_u)$, stacking real and imaginary parts, $\mathbf{W}^{(l)}$ being the weights at the l -th layer, $\mathbf{b}^{(l)}$ being the bias at the l -th layer, $\sigma(\cdot)$ as a leaky ReLU activation and the number of layers, L . This architecture ensures that $\mathbf{z}_u = \mathbf{h}^{(L)}$ captures symbol features while being resilient to channel distortions and interference.

After transmission through the multiuser channel, the received signal after MRC is $\mathbf{y}_u = [y_u(1), \dots, y_u(n)]^T$, where y_u is expressed as (17) where $r_k^{(u)}$ is given in (1) but now with z_u replacing s_u , which is the encoded latent representation. The MixDDPM denoiser $\mathcal{D}_\theta(\cdot)$, trained as described in Section IV-B, processes $\mathbf{y}_u$ to form a denoised latent vector

$$\hat{\mathbf{z}}_u = \mathcal{D}_\theta(\mathbf{y}_u), \quad (35)$$

in which θ parameterizes the reverse diffusion process. The denoiser leverages its mixture-based modeling to mitigate the multi-modal interference patterns unique to symbol-level FAMA, refining $\hat{\mathbf{z}}_u$ for subsequent decoding.

The JSCC decoder $\mathcal{D}_\psi(\cdot)$ reconstructs the original QPSK symbols from the denoised latent vector, i.e.,

$$\hat{\mathbf{s}}_u = \mathcal{D}_\psi(\hat{\mathbf{z}}_u), \quad (36)$$

where ψ represents the decoder's parameters. The decoder mirrors the encoder with transposed convolutional layers

$$\mathbf{h}^{(l)} = \sigma \left(\mathbf{W}^{(L-l+1)T} \mathbf{h}^{(l-1)} + \mathbf{b}^{(L-l+1)} \right), \text{ for } l = 1, \dots, L, \quad (37)$$

outputting $\hat{\mathbf{s}}_u$ in the complex domain by combining the final layer's real and imaginary components.

Our turbo FAMA employs a sequential, stage-wise training approach, first independently training the JSCC encoder-decoder and subsequently training the diffusion-based denoiser (MixDDPM). This separation inherently prevents gradient vanishing or instability issues caused by iterative denoising steps in the diffusion model from affecting the JSCC encoder-decoder. Thus, stable training and robust JSCC performance are preserved, even under challenging low-SNR conditions. Overall, the e2e system is trained in order to minimize the mean squared error between the transmitted and reconstructed symbols, augmented by the MixDDPM loss

$$\mathcal{L}_{\text{total}} = \mathbb{E}_{\mathbf{s}_u, \boldsymbol{\eta}_{\text{res}}} \left[\|\mathbf{s}_u - \mathcal{D}_\psi(\mathcal{D}_\theta(\mathcal{E}_\phi(\mathbf{s}_u) + \boldsymbol{\eta}_{\text{res}}))\|^2 \right] \quad (38)$$

$$+ \lambda \mathcal{L}_{\text{MixDDPM}},$$

$$\equiv \mathcal{L}_{\text{JSCC}} + \lambda \mathcal{L}_{\text{MixDDPM}}, \quad (39)$$

where λ balances reconstruction accuracy and denoising performance, and $\mathcal{L}_{\text{MixDDPM}}$ is defined in (29). The parameters ϕ , θ , and ψ are jointly optimized via gradient descent, enabling the system to adapt to the interference characteristics.

This integration enhances resilience of the proposed FAMA system by combining JSCC's ability to encode interference-robust representations with MixDDPM's capacity to remove structured residual interference. The training procedure for the combined JSCC and MixDDPM is outlined in Algorithm 4.

D. Model Structure

The proposed system employs a JSCC encoder that begins by embedding the input data into a latent representation through a series of patch-partition and convolutional layers, followed by multiple blocks integrating local convolutions and transformer-style self-attention to capture both short- and long-range dependencies. The encoded signal $\mathbf{z}$ is then transmitted over the channel, and the receiver applies a diffusion-based denoiser inspired by an improved U-Net architecture [102].

The corrupted signal $\mathbf{y}$ is passed to an initial convolutional layer, which prepares it for the U-Net pipeline [102]. Subsequent down-sampling blocks, via stride-based convolutions, reduce spatial resolution while increasing the feature-channel depth. Each U-Net stage consists of a convolutional residual (Conv-Res) block and a convolutional attention (Conv-Attn) block. The Conv-Res block replaces fully-connected layers in the residual path with additional convolutional layers, enhancing feature extraction at reduced computational cost. The Conv-Attn block applies a lightweight attention mechanism using depthwise or grouped convolutions to capture global dependencies. A timestep (or diffusion-step) embedding t is added in the intermediate layers to inform the network about the current stage of the reverse diffusion process.

After the deepest level of the U-Net, up-sampling blocks (interpolation plus a convolutional layer) gradually restore spatial resolution. Skip connections from the encoder stages are concatenated to the decoder stages, enabling the network to recover fine-grained details. The final U-Net output is passed through a projection layer that maps features back into the

Algorithm 4 Combined JSCC and MixDDPM Training

Require: Training set S , hyperparameters T, α_t, λ , CSI at the ports $\{g_k^{(u,u)}\}$, noise variance σ_η^2

Ensure: Trained JSCC encoder $\mathcal{E}_\phi$, decoder $\mathcal{D}_\psi$, and MixDDPM model p_θ

```

1: Stage 1: Pre-train JSCC without diffusion
2: while Training condition not met do
3:   Sample  $\mathbf{s} \sim S$ 
4:   Compute latent  $\mathbf{z} = \mathcal{E}_\phi(\mathbf{s})$ 
5:   Sample noise  $\mathbf{n} \sim \mathcal{N}(\mathbf{0}, \sigma_\eta^2 \mathbf{I})$ 
6:   Compute channel output  $\mathbf{y}$ 
7:   Update  $\phi, \psi$  to minimize  $\|\mathbf{s} - \mathcal{D}_\psi(\mathbf{y})\|^2$ 
8: end while
9: Stage 2: Train MixDDPM
10: while Training condition not met do
11:   Sample  $\mathbf{x}_0 \sim S$ 
12:   Sample timestep  $t \sim \text{Uniform}\{1, \dots, T\}$ 
13:   Sample noise  $\epsilon \sim \mathcal{N}(\mathbf{0}, \mathbf{I})$ 
14:   Compute  $\mathbf{x}_t = \sqrt{\alpha_t} \mathbf{x}_0 + \sqrt{1 - \alpha_t} \epsilon$ 
15:   Sample  $j \sim \text{Categorical}(\omega_1, \dots, \omega_K)$ 
16:   Compute loss:

```

$$\mathcal{L} = -\log p_\theta(j|\mathbf{x}_t) + \|\epsilon - \epsilon_\theta^{(j)}(\mathbf{x}_t)\|^2$$

```

17:   Update  $\theta$  using gradient descent on  $\mathcal{L}$ 
18: end while
19: Stage 3: Joint fine-tuning
20: while Training condition not met do
21:   Sample  $\mathbf{s} \sim S$ 
22:   Compute latent  $\mathbf{z} = \mathcal{E}_\phi(\mathbf{s})$ 
23:   Sample residual interference  $\eta_{\text{res}}$ 
24:   Compute received signal  $\mathbf{y}$  using (17)
25:   Denoise  $\hat{\mathbf{z}} = \mathcal{D}_\theta(\mathbf{y})$ 
26:   Reconstruct  $\hat{\mathbf{s}} = \mathcal{D}_\psi(\hat{\mathbf{z}})$ 
27:   Compute reconstruction loss  $\mathcal{L}_{\text{JSCC}} = \|\mathbf{s} - \hat{\mathbf{s}}\|^2$ 
28:   Compute MixDDPM loss  $\mathcal{L}_{\text{MixDDPM}}$  using  $\theta$ 
29:   Update  $\phi, \theta, \psi$  to minimize  $\mathcal{L}_{\text{JSCC}} + \lambda \mathcal{L}_{\text{MixDDPM}}$ 
30: end while

```

symbol domain. This denoised signal is then fed to a JSCC decoder, mirroring the encoder's structure but replacing down-sampling modules with the corresponding up-sampling modules to reconstruct the original data. By blending convolution-based feature refinement, attention-driven global modeling, and iterative diffusion, this architecture effectively learns to remove non-Gaussian residual interference characteristics.

Note also that while we have assumed perfect CSI at UTs for simplicity, practical deployments will inevitably experience CSI estimation errors due to pilot contamination and limited training. Such errors can degrade port shortlisting accuracy by affecting the evaluation of desired-signal power levels. However, our heuristic shortlisting method inherently exhibits tolerance to moderate CSI inaccuracies. Moreover, the proposed MixDDPM diffusion-based denoising approach compensates for such imperfections by learning from data-driven interference distributions, effectively reducing residual interference even under suboptimal port selection.

TABLE I
SIMULATION PARAMETERS

Parameter	Value
T (training/sampling)	1000/50
mixture weights, ω	[0.4, 0.15, 0.15, 0.1, 0.1, 0.05, 0.05]
mixture centres, c	[0, 1, -1, $\sqrt{2}$, $-\sqrt{2}$, $\sqrt{3}$, $-\sqrt{3}$]
α_t	0.9999 to 0.9800
Optimizer	ADAM
Learning rate	0.0001
dropout	0.1
Residual blocks	2
Embedding dimensions (enc.)	[128, 256]
Patch size (enc.)	2
Input channels (enc.)	2
Attention depth (enc.)	[6, 6]
Attention heads (enc.)	[8, 16]
Window size (enc.)	4
Embedding dimensions (dec.)	[256, 128]
Attention depth (dec.)	[6, 6]
Attention heads (dec.)	[16, 8]
Window size (dec.)	4
Batch size	64
Training epochs (JSCC)	100
Training epochs (MixDDPM)	100
Training epochs (Joint)	20

The overall proposed approach is referred to as turbo FAMA and Fig. 3 illustrates the blocks of the e2e system.

V. SIMULATION RESULTS

In this section, we evaluate the performance of the proposed turbo FAMA system under different settings. We also include two important benchmarks for comparison. All of the methods considered are explained below:

- Turbo FAMA (with JSCC)—This is the proposed technique that adopts JSCC, port shortlisting, and MRC.
- Turbo FAMA (with JSCC+MixDDPM)—This is the same as above plus the use of MixDDPM after MRC.
- All-port MRC—This scheme utilizes all the port received signals for MRC to produce the output signal. Both JSCC and MixDDPM are also used for IUI immunity.
- Fast FAMA—This scheme assumes perfect knowledge of the port instantaneous SINRs at each UT to find the best port maximizing it. Also, we consider an excessively large number of ports just for this scheme with $K = 1000$. This represents the ideal case for fast FAMA.

Simulation results are conducted under many different settings by varying the number of users, U , the number of ports, K , and the normalized length of FAS, W . Rich scattering Rayleigh fading channels are considered. The average SNR is set to 10 dB and the blocklength before the encoder is 1024. We also define the channel bandwidth ratio, $\text{CBR} \triangleq \frac{n}{1024}$, to measure the bandwidth expansion when utilizing JSCC. If $n = 1024$, then no bandwidth expansion is introduced. In the simulations, unless specified otherwise, $\text{CBR} = 1$. Also, in the port shortlisting process, we set $K_{\text{sel}} = 20$ and $\gamma_{\text{th}} = 0.6$. Note that the parameters for the neural network structures used for JSCC and MixDDPM are listed in TABLE I.

Results in Fig. 4 are provided for the SER performance for different approaches against the number of users sharing the same physical channel, from a few users to as large as 1000 users. Note that all the approaches do not require CSI at the

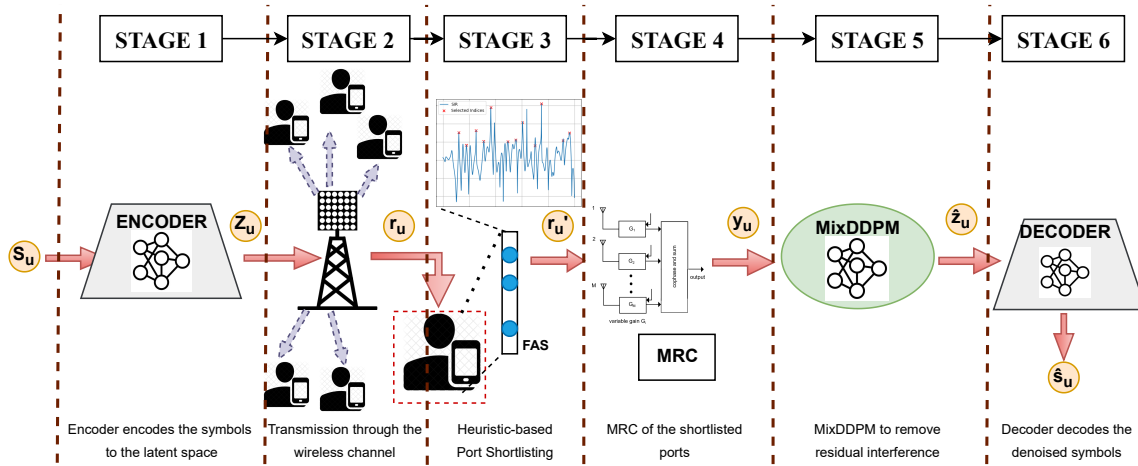


Fig. 3. The proposed e2e turbo FAMA system with port shortlisting, MRC, MixDDPM and JSCC.

BS and no precoding is applied whatsoever. The number of ports at the FAS for each UT is set as $K = 200$ except for the fast FAMA method and $W = 20$ is assumed. The results illustrate that as the number of users increases, the SER of all the schemes degrades due to elevated IUI, which is expected. Additionally, the SER initially increases sharply, reflecting the challenges in distinguishing desired signals from IUI as the number of users rises in massive scenarios, before gradually plateauing. The results also show that all-port MRC does not really work unless the number of users is small. The ideal fast FAMA scheme does better but is unable to handle more than 200 users. By contrast, the proposed turbo FAMA approaches outperform both all-port MRC and fast FAMA significantly. The use of MixDDPM in turbo FAMA seems to be particularly effective when the number of users is not so large (say ≤ 200 users) but its effectiveness diminishes if the number of users becomes really large⁶. Remarkably, the results reveal that the proposed turbo FAMA scheme can handle up to 200 users if $\text{SER} = 0.01$ is required. If a higher SER, say $\text{SER} = 0.1$, is acceptable, then turbo FAMA can even deal with 1000 users, all without the need of precoding at the BS.

In Fig. 5, we investigate the impact of the number of ports, K , on the SER performance for all these FAMA schemes. In the results, we have fixed $U = 200$ users while the number of ports is changed. Evidently, we expect that as K increases, the general trend should be that SER decreases. This is clearly what happens to fast FAMA because a larger K will give fast FAMA a higher spatial resolution to avoid IUI. Nevertheless, for the other three methods, all-port MRC and the two turbo FAMA schemes, the SER first drops and then saturates when K reaches a certain point. For all-port MRC, it is clear that higher spatial resolution does not necessarily imply better IUI immunity as MRC does not take into account any knowledge of IUI when combining. For turbo FAMA, on the other hand, it may be explained by the fact that the performance gain of port shortlisting diminishes as K becomes large, meaning that the shortlisted ports do not change much if K continues to

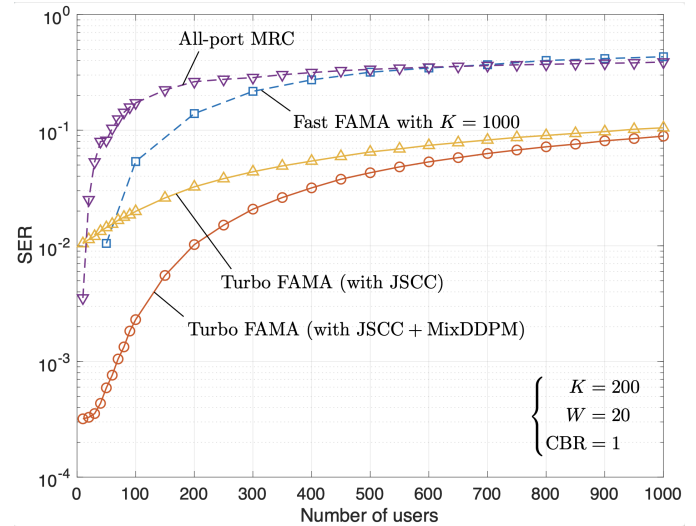


Fig. 4. SER versus the number of users for different FAMA approaches with $K = 200$, $W = 20$, and $\text{CBR} = 1$.

increase beyond a certain value. This is in fact a welcoming property because there appears to have an optimal value of K to keep the complexity of training and hardware of FAS low when near-optimal SER performance is achieved.

Fig. 6 illustrates the SER performance versus the normalized size of the fluid antenna at each UT, W . As expected, when W increases, the spatial diversity is enhanced because a larger antenna aperture provides more spatial samples, leading to a significant reduction in SER. Moreover, the FAS size affects fast FAMA more than other schemes. In particular, the results suggest that the SER of turbo FAMA saturates more quickly as W increases. This may be explained by the fact that the SDM approach in the port shortlisting process already decorrelates the channels of the shortlisted ports very well even when W is small and thus there is a diminishing return when W becomes large. Additionally, the results demonstrate that MixDDPM is a lot more effective for large W . This means that the diffusion-based denoiser relies on the FAS size to work well.

Now, we turn our attention to study the importance of SDM

⁶This performance can be improved by integrating with other schemes specifically designed for strong interference scenarios, such as in [105].

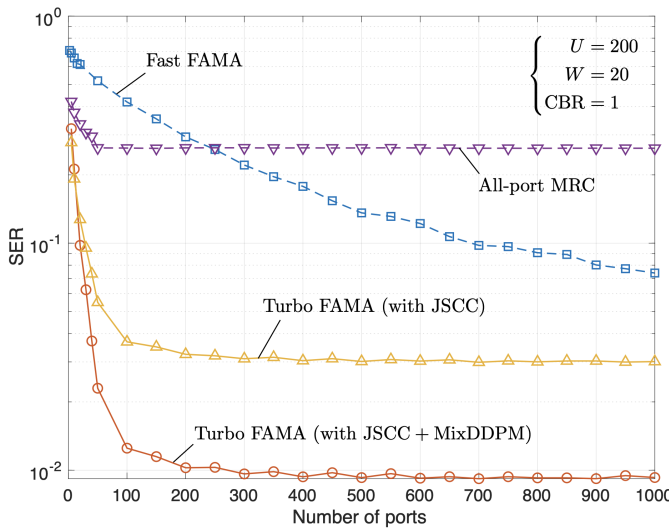


Fig. 5. SER versus the number of ports for different FAMA approaches with $U = 200$, $W = 20$, and $CBR = 1$.

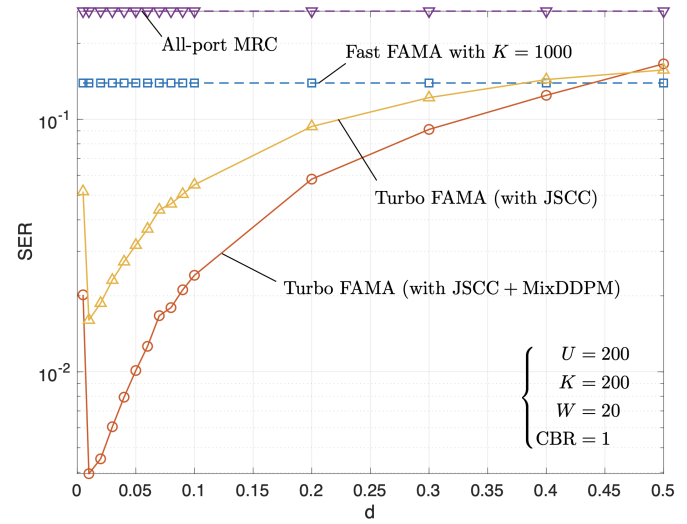


Fig. 7. SER versus the spacing between shortlisted ports for different FAMA approaches with $K = 200$, $U = 200$, $W = 20$ and $CBR = 1$.

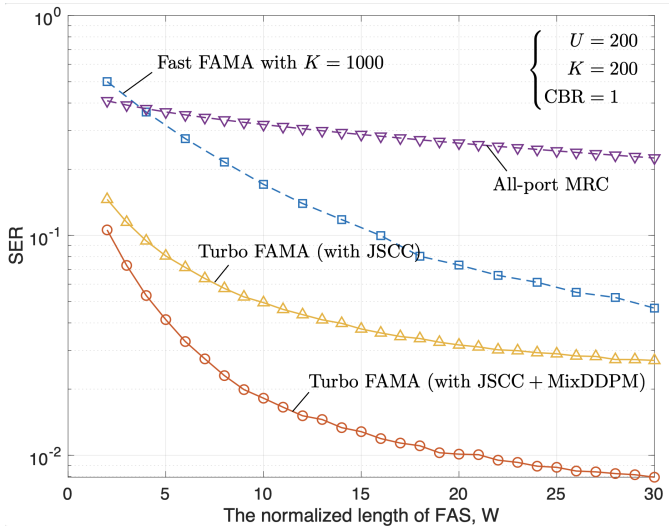


Fig. 6. SER versus the normalized size of FAS at each UT for different FAMA approaches with $K = 200$, $U = 200$, and $CBR = 1$.

in the port shortlisting process. To do so, in the turbo FAMA schemes, we disable the SDM step but instead enforce a fixed minimum spacing of $d \times K$ between the shortlisted ports from the set $\mathcal{K}^{(u)}$. The parameter, d , allows us to control the spacing between the shortlisted ports and then highlight the importance in the spatial diversity of the shortlisted ports. Note that the spacing parameter, d , only affects the proposed turbo FAMA schemes but not all-port MRC and fast FAMA. The corresponding SER results are given in Fig. 7. As we can see, the spacing plays an important role in the SER performance for the turbo FAMA schemes. Specifically, when d increases, the spatial correlation between the shortlisted ports decreases, leading to more independent channel realizations that should facilitate more effective interference mitigation. However, this reduces the possible number of shortlisted ports that degrades the overall performance. This is why for very small values of d , although the ports are closely spaced and more correlated, the SER performance is still better. While this may suggest that

the smaller the value of d , the better the SER performance, it is important to recognize that when d becomes too small, the SER shoots up. This abrupt rise in the SER can be explained by the fact that when d is too small, there are too many closely spaced shortlisted ports with low diversity, and after MRC, gives a degraded signal in terms of IUI. Therefore, there is an optimal value of d which is small but should not be too small for minimizing the SER. For example, the results show that turbo FAMA with MixDDPM (but without SDM) at $d = 0.05$ obtains $SER = 0.01$, the level achieved by turbo FAMA with MxDDPM and SDM. If $d = 0.01$, then turbo FAMA achieves the minimum SER of 0.004. However, finding the optimal d is tedious and requires retraining of JSCC and MxDDPM. The use of SDM thus can be viewed as a smart compromise, which balances between diversity and the number of shortlisted ports, and eliminates the need of choosing d .

When using JSCC, it is actually possible to either introduce redundancy for improved error correction or compression for bandwidth saving. This effect can be considered by considering different values of CBR. Note that JSCC is employed for all-port MRC and the two turbo FAMA schemes except the fast FAMA method. The results in Fig. 8 illustrate that the SER decreases as the CBR increases, which will make all-port MRC approach the performance of fast FAMA. When the CBR increases, the JSCC encoder utilizes additional redundancy to enhance protection against interference-induced distortions. This redundancy is intelligently embedded into the transmitted representation through learned mappings optimized directly from the data. As the redundancy grows, the deep JSCC encoder and decoder jointly learn richer latent representations, enabling better recovery from severe interference events and effectively reducing SER. Additionally, it is possible to achieve bandwidth saving using JSCC in the proposed turbo FAMA schemes with some mild SER degradation.

We next provide results to investigate how SER changes according to the blocklength of the data inputting to the JSCC encoder. As before, these results only affect all-port MRC and

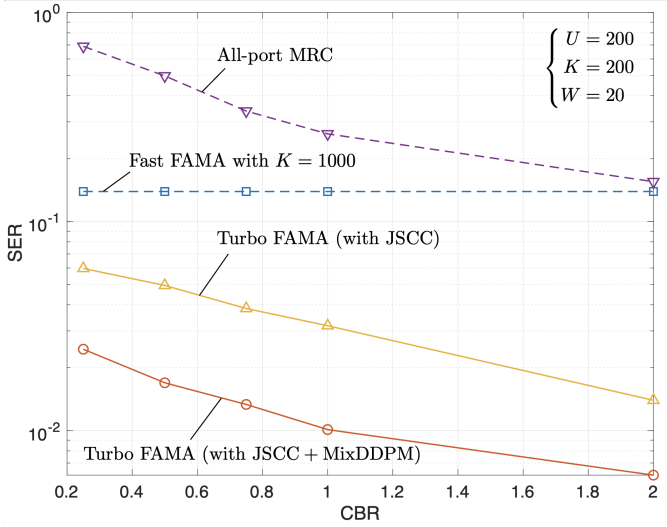


Fig. 8. SER versus the CBR for different FAMA approaches with $K = 200$, $U = 200$, and $W = 20$.

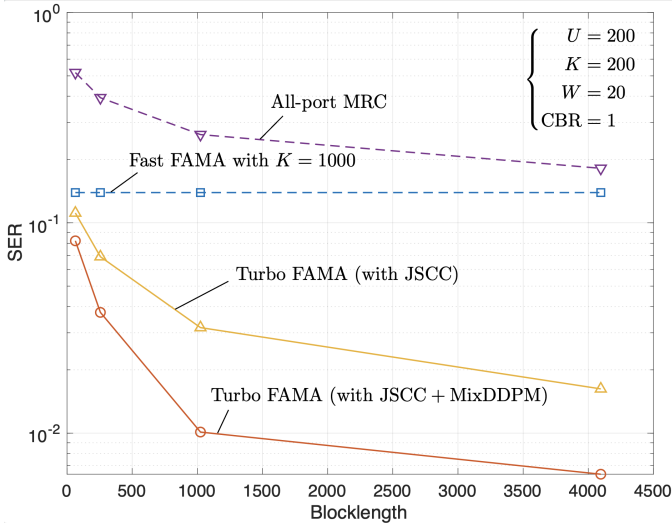


Fig. 9. SER versus the blocklength for different FAMA approaches with $K = 200$, $U = 200$, $W = 20$ and $\text{CBR} = 1$.

the two turbo FAMA schemes as fast FAMA does not apply JSCC. These results are provided in Fig. 9, showing that as the blocklength increases, all schemes exhibit improved SER. However, beyond a certain point, the SER improvement tends to saturate. This saturation can be attributed to diminishing returns in redundancy gains without increasing CBR, as well as limitations in training complexity.

Finally, we evaluate the robustness of the proposed turbo FAMA scheme under practical CSI estimation errors at UTs. The imperfect CSI at each UT is modeled as

$$\hat{g}_k^{(u,u)} = g_k^{(u,u)} + e_k^{(u,u)}, \quad (40)$$

in which $g_k^{(u,u)}$ denotes the true channel gain, and $e_k^{(u,u)} \sim \mathcal{CN}(0, \sigma_e^2)$ represents the Gaussian CSI errors with variance σ_e^2 . Fig. 10 illustrates the SER results against the standard deviation of the CSI errors, σ_e , for different schemes. Ideal fast FAMA maintains constant performance due to the assumption of having perfect CSI. The proposed MixDDPM+JSCC shows

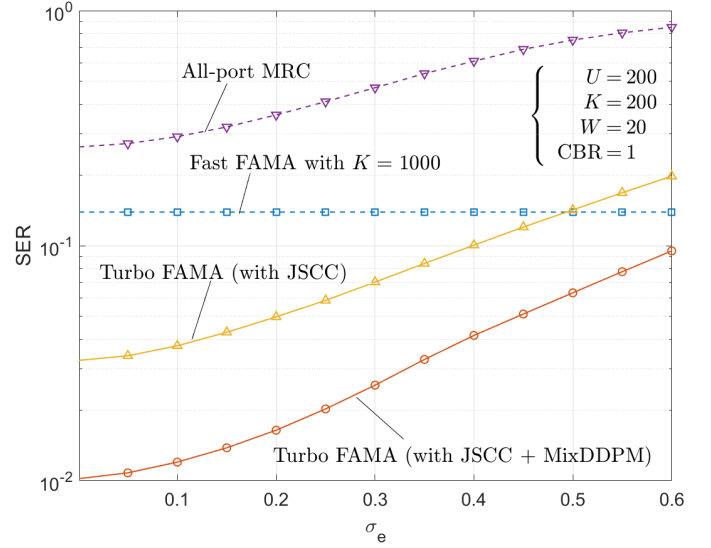


Fig. 10. SER versus the standard deviation of CSI errors (σ_e) for different FAMA approaches with $K = 200$, $U = 200$, $W = 20$ and $\text{CBR} = 1$.

strong resilience with minor SER degradation, owing to its implicit denoising capability. However, JSCC-only appears to be more sensitive, reflecting its higher sensitivity to incorrect port selection stemming from CSI errors, and the conventional MRC suffers significant degradation due to heavy reliance on accurate CSI. These results highlight turbo FAMA's effectiveness in realistic scenarios with imperfect CSI.

VI. CONCLUSION

In this paper, we proposed a new FAMA scheme that utilizes position flexibility in FAS at UTs to realize extreme massive access, without the need for precoding at the BS. The proposed scheme is dubbed turbo FAMA as it integrates heuristic port shortlisting, MRC, JSCC and diffusion-based denoising using MixDDPM into one. The overall turbo FAMA system is able to shortlist promising ports for obtaining favourable received signals for MRC, which is followed by MixDDPM denoising to suppress non-Gaussian residual interference, and then JSCC decoding. Our simulation results point to a promising future, showing that it is possible to support 200 users on the same channel using turbo FAMA if $\text{SER} = 0.01$ is required. Further, if only $\text{SER} = 0.1$ is needed, the proposed turbo FAMA can even serve 1000 users, truly answering the scalability issue in massive access scenarios. Future work should look at ways to reduce the complexity of obtaining the received signals at all the FAS ports, and study the performance impact if only some of the received signals at the FAS are available.

REFERENCES

- [1] J. G. Andrews, F. Baccelli and R. K. Ganti, "A tractable approach to coverage and rate in cellular networks," *IEEE Trans. Commun.*, vol. 59, no. 11, pp. 3122–3134, Nov. 2011.
- [2] Z. Zhang *et al.*, "6G wireless networks: Vision, requirements, architecture, and key technologies," *IEEE Veh. Technol. Mag.*, vol. 14, no. 3, pp. 28–41, Sept. 2019.
- [3] W. Saad, M. Bennis, and M. Chen, "A vision of 6G wireless systems: Applications, trends, technologies, and open research problems," *IEEE Netw.*, vol. 34, no. 3, pp. 134–142, May/Jun. 2020.

- [4] F. Tariq *et al.*, "A speculative study on 6G," *IEEE Wireless Commun.*, vol. 27, no. 4, pp. 118–125, Aug. 2020.
- [5] ETSI, "ETSI's group on multiple access techniques for 6G networks," Dec. 2023.
- [6] K.-K. Wong, R. D. Murch, R.-K. Cheng, and K. B. Letaief, "Optimizing the spectral efficiency of multiuser MIMO smart antenna systems," in *Proc. IEEE Wireless Commun. Netw. Conf.*, vol. 1, pp. 426–430, 23–28 Sept. 2000, Chicago, IL, USA.
- [7] K.-K. Wong, R. D. Murch and K. B. Letaief, "Performance enhancement of multiuser MIMO wireless communication systems," *IEEE Trans. Commun.*, vol. 50, no. 12, pp. 1960–1970, Dec. 2002.
- [8] K. K. Wong, R. D. Murch, and K. Ben Letaief, "A joint-channel diagonalization for multiuser MIMO antenna systems," *IEEE Trans. Wireless Commun.*, vol. 4, no. 2, pp. 773–786, Jul. 2003.
- [9] S. Vishwanath, N. Jindal and A. Goldsmith, "Duality, achievable rates, and sum-rate capacity of Gaussian MIMO broadcast channels," *IEEE Trans. Inf. Theory*, vol. 49, no. 10, pp. 2658–2668, Oct. 2003.
- [10] L.-U. Choi and R. D. Murch, "A transmit preprocessing technique for multiuser MIMO systems using a decomposition approach," *IEEE Trans. Wireless Commun.*, vol. 3, no. 1, pp. 20–24, Jan. 2004.
- [11] Q. H. Spencer, A. L. Swindlehurst and M. Haardt, "Zero-forcing methods for downlink spatial multiplexing in multiuser MIMO channels," *IEEE Trans. Sig. Process.*, vol. 52, no. 2, pp. 461–471, Feb. 2004.
- [12] T. L. Marzetta, "Noncooperative cellular wireless with unlimited numbers of base station antennas," *IEEE Trans. Wireless Commun.*, vol. 9, no. 11, pp. 3590–3600, Nov. 2010.
- [13] E. G. Larsson, O. Edfors, F. Tufvesson and T. L. Marzetta, "Massive MIMO for next generation wireless systems," *IEEE Commun. Mag.*, vol. 52, no. 2, pp. 186–195, Feb. 2014.
- [14] D. A. Urquiza Villalonga, H. OdetAlla, M. J. Fernández-Getino Garcia, and A. Flizkowski, "Spectral efficiency of precoded 5G-NR in single and multi-user scenarios under imperfect channel knowledge: A comprehensive guide for implementation," *Electronics*, vol. 11, no. 24 (article number: 4237), Dec. 2022.
- [15] Z. Wang *et al.*, "Extremely large-scale MIMO: Fundamentals, challenges, solutions, and future directions," *IEEE Wireless Commun.*, vol. 31, no. 3, pp. 117–124, Jun. 2024.
- [16] H. Q. Ngo, A. Ashikhmin, H. Yang, E. G. Larsson and T. L. Marzetta, "Cell-free massive MIMO versus small cells," *IEEE Trans. Wireless Commun.*, vol. 16, no. 3, pp. 1834–1850, Mar. 2017.
- [17] Y. Liu, C. Ouyang, Z. Ding, and R. Schober, "The road to next-generation multiple access: A 50-year tutorial review," *arXiv preprint, arXiv:2403.00189*, 2024.
- [18] Y. Saito *et al.*, "Non-orthogonal multiple access (NOMA) for cellular future radio access," in *Proc. IEEE Veh. Technol. Conf. (VTC Spring)*, 2–5 Jun. 2013, Dresden, Germany.
- [19] P. Wang, J. Xiao and L. Ping, "Comparison of orthogonal and non-orthogonal approaches to future wireless cellular systems," *IEEE Veh. Technol. Mag.*, vol. 1, no. 3, pp. 4–11, Sept. 2006.
- [20] L. Dai *et al.*, "Non-orthogonal multiple access for 5G: Solutions, challenges, opportunities, and future research trends," *IEEE Commun. Mag.*, vol. 53, no. 9, pp. 74–81, Sept. 2015.
- [21] Y. Liu *et al.*, "Evolution of NOMA toward next generation multiple access (NGMA) for 6G," *IEEE J. Select. Areas Commun.*, vol. 40, no. 4, pp. 1037–1071, Apr. 2022.
- [22] Y. Mao *et al.*, "Rate-splitting multiple access: Fundamentals, survey, and future research trends," *IEEE Commun. Surv. & Tut.*, vol. 24, no. 4, pp. 2073–2126, Fourthquarter 2022.
- [23] B. Clerckx *et al.*, "A primer on rate-splitting multiple access: Tutorial, myths, and frequently asked questions," *IEEE J. Select. Areas Commun.*, vol. 41, no. 5, pp. 1265–1308, May 2023.
- [24] K. K. Wong and K. F. Tong, "Fluid antenna multiple access," *IEEE Trans. Wireless Commun.*, vol. 21, no. 7, pp. 4801–4815, Jul. 2022.
- [25] K.-K. Wong, K.-F. Tong, Y. Chen, and Y. Zhang, "Fast fluid antenna multiple access enabling massive connectivity," *IEEE Commun. Lett.*, vol. 27, no. 2, pp. 711–715, Feb. 2023.
- [26] K. K. Wong, K. F. Tong, Y. Shen, Y. Chen, and Y. Zhang, "Bruce Lee-inspired fluid antenna system: Six research topics and the potentials for 6G," *Frontiers Commun. and Netw., section Wireless Commun.*, vol. 3, no. 853416, Mar. 2022.
- [27] K. K. Wong, W. K. New, X. Hao, K. F. Tong, and C. B. Chae, "Fluid antenna system—Part I: Preliminaries," *IEEE Commun. Lett.*, vol. 27, no. 8, pp. 1919–1923, Aug. 2023.
- [28] W. K. New *et al.*, "A tutorial on fluid antenna system for 6G networks: Encompassing communication theory, optimization methods and hardware designs," *IEEE Commun. Surv. & Tut.*, vol. 24, no. 4, pp. 2325–2377, Aug. 2025.
- [29] K. K. Wong, A. Shojaeifard, K.-F. Tong, and Y. Zhang, "Performance limits of fluid antenna systems," *IEEE Commun. Lett.*, vol. 24, no. 11, pp. 2469–2472, Nov. 2020.
- [30] K. K. Wong, A. Shojaeifard, K.-F. Tong, and Y. Zhang, "Fluid antenna systems," *IEEE Trans. Wireless Commun.*, vol. 20, no. 3, pp. 1950–1962, Mar. 2021.
- [31] W.-J. Lu *et al.*, "Fluid antennas: Reshaping intrinsic properties for flexible radiation characteristics in intelligent wireless networks," accepted in *IEEE Commun. Mag.*, vol. 63, no. 5, pp. 40–45, May 2025.
- [32] J. Zhang *et al.*, "A novel pixel-based reconfigurable antenna applied in fluid antenna systems with high switching speed," *IEEE Open J. Antennas & Propag.*, vol. 6, no. 1, pp. 212–228, Feb. 2025.
- [33] C. Wang *et al.*, "AI-empowered fluid antenna systems: Opportunities, challenges and future directions," *IEEE Wireless Commun.*, vol. 31, no. 5, pp. 34–41, Oct. 2024.
- [34] Y. Huang, L. Xing, C. Song, S. Wang, and F. Elhouni, "Liquid antennas: Past, present and future," *IEEE Open J. Antennas & Propag.*, vol. 2, pp. 473–487, Mar. 2021.
- [35] Y. Shen *et al.*, "Design and implementation of mmWave surface wave enabled fluid antennas and experimental results for fluid antenna multiple access," *arXiv preprint, arXiv:2405.09663*, May 2024.
- [36] R. Wang *et al.*, "Electromagnetically reconfigurable fluid antenna system for wireless communications: Design, modeling, algorithm, fabrication, and experiment," *arXiv preprint, arXiv:2502.19643v2*, 2025.
- [37] S. Basbug, "Design and synthesis of antenna array with movable elements along semicircular paths," *IEEE Antennas Wireless Propag. Lett.*, vol. 16, pp. 3059–3062, Oct. 2017.
- [38] B. Liu, K.-F. Tong, K. K. Wong, C.-B. Chae, and H. Wong, "Programmable meta-fluid antenna for spatial multiplexing in fast fluctuating radio channels," *Optics Express*, vol. 33, no. 13, pp. 28898–28915, 2025.
- [39] S. Shen, K. K. Wong, and R. Murch, "Antenna coding empowered by pixel antennas," *arXiv preprint, arXiv:2411.06642*, 2024.
- [40] M. Khammassi, A. Kammoun and M.-S. Alouini, "A new analytical approximation of the fluid antenna system channel," *IEEE Trans. Wireless Commun.*, vol. 22, no. 12, pp. 8843–8858, Dec. 2023.
- [41] W. K. New, K. K. Wong, H. Xu, K. F. Tong and C.-B. Chae, "Fluid antenna system: New insights on outage probability and diversity gain," *IEEE Trans. Wireless Commun.*, vol. 23, no. 1, pp. 128–140, Jan. 2024.
- [42] P. Ramírez-Espinosa, D. Morales-Jimenez, and K. K. Wong, "A new spatial block-correlation model for fluid antenna systems," *IEEE Trans. Wireless Commun.*, vol. 23, no. 11, pp. 15829–15843, Nov. 2024.
- [43] J. D. Vega-Sánchez, L. Urquiza-Aguiar, M. C. P. Paredes, and D. P. M. Osorio, "A simple method for the performance analysis of fluid antenna systems under correlated Nakagami- m fading," *IEEE Wireless Commun. Lett.*, vol. 13, no. 2, pp. 377–381, Feb. 2024.
- [44] J. D. Vega-Sánchez, A. E. López-Ramírez, L. Urquiza-Aguiar, and D. P. M. Osorio, "Novel expressions for the outage probability and diversity gains in fluid antenna system," *IEEE Wireless Commun. Lett.*, vol. 13, no. 2, pp. 372–376, Feb. 2024.
- [45] P. D. Alvim *et al.*, "On the performance of fluid antennas systems under α - μ fading channels," *IEEE Wireless Commun. Lett.*, vol. 13, no. 1, pp. 108–112, Jan. 2024.
- [46] F. Rostami Ghadi, K. K. Wong, F. Javier López-Martínez, and K. F. Tong, "Copula-based performance analysis for fluid antenna systems under arbitrary fading channels," *IEEE Commun. Lett.*, vol. 27, no. 11, pp. 3068–3072, Nov. 2023.
- [47] F. Rostami Ghadi *et al.*, "A Gaussian copula approach to the performance analysis of fluid antenna systems," *IEEE Trans. Wireless Commun.*, vol. 23, no. 11, pp. 17573–17585, Nov. 2024.
- [48] W. K. New, K. K. Wong, H. Xu, K. F. Tong, and C.-B. Chae, "An information-theoretic characterization of MIMO-FAS: Optimization, diversity-multiplexing tradeoff and q -outage capacity," *IEEE Trans. Wireless Commun.*, vol. 23, no. 6, pp. 5541–5556, Jun. 2024.
- [49] H. Xu *et al.*, "Channel estimation for FAS-assisted multiuser mmWave systems," *IEEE Commun. Lett.*, vol. 28, no. 3, pp. 632–636, Mar. 2024.
- [50] Z. Zhang, J. Zhu, L. Dai, and R. W. Heath Jr, "Successive Bayesian reconstructor for channel estimation in fluid antenna systems," *IEEE Trans. Wireless Commun.*, vol. 24, no. 3, pp. 1992–2006, Mar. 2025.
- [51] B. Xu, Y. Chen, Q. Cui, X. Tao, and K. K. Wong, "Sparse Bayesian learning-based channel estimation for fluid antenna systems," *IEEE Wireless Commun. Lett.*, vol. 14, no. 2, pp. 325–329, Feb. 2025.
- [52] W. K. New *et al.*, "Channel estimation and reconstruction in fluid antenna system: Oversampling is essential," *IEEE Trans. Wireless Commun.*, vol. 24, no. 1, pp. 309–322, Jan. 2025.
- [53] B. Tang *et al.*, "Fluid antenna enabling secret communications," *IEEE Commun. Lett.*, vol. 27, no. 6, pp. 1491–1495, Jun. 2023.

- [54] H. Xu *et al.*, "Coding-enhanced cooperative jamming for secret communication in fluid antenna systems," *IEEE Commun. Lett.*, vol. 28, no. 9, pp. 1991–1995, Sept. 2024.
- [55] F. R. Ghadi *et al.*, "Physical layer security over fluid antenna systems: Secrecy performance analysis," *IEEE Trans. Wireless Commun.*, vol. 23, no. 12, pp. 18201–18213, Dec. 2024.
- [56] C. Skouroumounis and I. Krikidis, "Fluid antenna-aided full duplex communications: A macroscopic point-of-view," *IEEE J. Select. Areas Commun.*, vol. 41, no. 9, p. 2879–2892, Sept. 2023.
- [57] F. Rostami Ghadi *et al.*, "On performance of RIS-aided fluid antenna systems," *IEEE Wireless Commun. Lett.*, vol. 13, no. 8, pp. 2175–2179, Aug. 2024.
- [58] A. Salem, K.-K. Wong, G. Alexandropoulos, C.-B. Chae, and R. Murch, "A first look at the performance enhancement potential of fluid reconfigurable intelligent surface," *arXiv preprint, arXiv:2502.17116*, 2025.
- [59] H. Yang *et al.*, "Position index modulation for fluid antenna system," *IEEE Trans. Wireless Commun.*, vol. 23, no. 11, pp. 16773–16787, Nov. 2024.
- [60] J. Zhu *et al.*, "Fluid antenna empowered index modulation for RIS-aided mmWave transmissions," *IEEE Trans. Wireless Commun.*, vol. 24, no. 2, pp. 1635–1647, Feb. 2025.
- [61] C. Wang *et al.*, "Fluid antenna system liberating multiuser MIMO for ISAC via deep reinforcement learning," *IEEE Trans. Wireless Commun.*, vol. 23, no. 9, pp. 10879–10894, Sept. 2024.
- [62] L. Zhou, J. Yao, M. Jin, T. Wu and K. K. Wong, "Fluid antenna-assisted ISAC systems," *IEEE Wireless Commun. Lett.*, vol. 13, no. 12, pp. 3533–3537, Dec. 2024.
- [63] J. Zou *et al.*, "Shifting the ISAC trade-off with fluid antenna systems," *IEEE Wireless Commun. Lett.*, vol. 13, no. 12, pp. 3479–3483, Dec. 2024.
- [64] K. Meng, C. Masouros, K. K. Wong, A. P. Petropulu, and L. Hanzo, "Integrated sensing and communication meets smart propagation engineering: Opportunities and challenges," *IEEE Netw.*, vol. 39, no. 2, pp. 278–285, Mar. 2025.
- [65] L. Zhu and K. K. Wong, "Historical review of fluid antenna and movable antenna," *arXiv preprint, arXiv:2401.02362v2*, 2024.
- [66] J. Zheng *et al.*, "Flexible-position MIMO for wireless communications: Fundamentals, challenges, and future directions," *IEEE Wireless Commun.*, vol. 31, no. 5, pp. 18–26, Oct. 2024.
- [67] Z. Yang *et al.*, "Pinching antennas: Principles, applications and challenges," *arXiv preprint, arXiv:2501.10753*, 2025.
- [68] K. K. Wong, D. Morales-Jimenez, K. F. Tong, and C. B. Chae, "Slow fluid antenna multiple access," *IEEE Trans. Commun.*, vol. 71, no. 5, pp. 2831–2846, May 2023.
- [69] H. Xu *et al.*, "Revisiting outage probability analysis for two-user fluid antenna multiple access system," *IEEE Trans. Wireless Commun.*, vol. 23, no. 8, pp. 9534–9548, Aug. 2024.
- [70] N. Waqar, K. K. Wong, K. F. Tong, A. Sharples, and Y. Zhang, "Deep learning enabled slow fluid antenna multiple access," *IEEE Commun. Lett.*, vol. 27, no. 3, pp. 861–865, Mar. 2023.
- [71] M. Eskandari, A. Burr, K. Cumanan, and K. K. Wong, "cGAN-based slow fluid antenna multiple access," *IEEE Wireless Commun. Lett.*, vol. 13, no. 10, pp. 2907–2911, Oct. 2024.
- [72] K. K. Wong, C. B. Chae, and K. F. Tong, "Compact ultra massive antenna array: A simple open-loop massive connectivity scheme," *IEEE Trans. Wireless Commun.*, vol. 23, no. 6, pp. 6279–6294, Jun. 2024.
- [73] K. K. Wong, "Transmitter CSI-free RIS-randomized CUMA for extreme massive connectivity," *IEEE Open J. Commun. Soc.*, vol. 5, pp. 6890–6902, 2024.
- [74] H. Hong, K. K. Wong, K. F. Tong, H. Shin, and Y. Zhang, "Coded fluid antenna multiple access over fast fading channels," *IEEE Wireless Commun. Lett.*, vol. 14, no. 4, pp. 1249–1253, Apr. 2025.
- [75] H. Hong *et al.*, "Downlink OFDM-FAMA in 5G-NR systems," *IEEE Trans. Wireless Commun.*, doi:10.1109/TWC.2025.3577771, 2025.
- [76] N. Waqar *et al.*, "Opportunistic fluid antenna multiple access via team-inspired reinforcement learning," *IEEE Trans. Wireless Commun.*, vol. 23, no. 9, pp. 12068–12083, Sept. 2024.
- [77] W. K. New *et al.*, "Fluid antenna system enhancing orthogonal and non-orthogonal multiple access," *IEEE Commun. Letters*, vol. 28, no. 1, pp. 218–222, Jan. 2024.
- [78] F. Rostami Ghadi, K. K. Wong, F. J. López-Martínez, L. Hanzo, and C.-B. Chae, "Fluid antenna-aided rate-splitting multiple access," *IEEE Trans. Veh. Technol.*, doi:10.1109/TVT.2025.3599709, 2025.
- [79] A. Zappone, M. Di Renzo, and M. Debbah, "Wireless networks design in the era of deep learning: Model-based, AI-based, or both?" *IEEE Trans. Commun.*, vol. 67, no. 10, pp. 7331–7376, Oct. 2019.
- [80] U. Challita, W. Saad, and C. Bettstetter, "Interference management for cellular-connected UAVs: A deep reinforcement learning approach," *IEEE Trans. Wireless Commun.*, vol. 18, no. 4, pp. 2125–2140, Apr. 2019.
- [81] Y. Shen, Y. Shi, J. Zhang, and K. B. Letaief, "Graph neural networks for scalable radio resource management: Architecture design and theoretical analysis," *IEEE J. Select. Areas Commun.*, vol. 39, no. 1, pp. 101–115, Jan. 2021.
- [82] W. Zhang, J. Yin, D. Wu, G. Guo and Z. Lai, "A self-interference cancellation method based on deep learning for beyond 5G full-duplex system," in *Proc. IEEE Int. Conf. Sig. Process., Commun. & Comput. (ICSPCC)*, 14–16 Sept. 2018, Qingdao, China.
- [83] P. Lohan, B. Kantarci, M. A. Ferrag, N. Tihanyi, and Y. Shi, "From 5G to 6G networks: A survey on AI-based jamming and interference detection and mitigation," *IEEE Open J. Commun. Soc.*, vol. 5, pp. 3920–3974, 2024.
- [84] Q. Wu and R. Zhang, "Towards smart and reconfigurable environment: Intelligent reflecting surface aided wireless network," *IEEE Commun. Mag.*, vol. 58, no. 1, pp. 106–112, Jan. 2020.
- [85] N. Waqar *et al.*, "Computation offloading and resource allocation in MEC-enabled integrated aerial-terrestrial vehicular networks: A reinforcement learning approach," *IEEE Trans. Intelligent Transport. Syst.*, vol. 23, no. 11, pp. 21478–21491, Nov. 2022.
- [86] M. Letafati, S. Ali, and M. Latva-aho, "Diffusion models for wireless communications," *arXiv preprint, arXiv:2310.07312v3*, 2023.
- [87] X. Wang, P. Zheng, and N. Cheng, "Erasing noise in signal detection with diffusion model: From theory to application," *arXiv preprint, arXiv:2501.07030*, 2025.
- [88] H. Du *et al.*, "Enhancing deep reinforcement learning: A tutorial on generative diffusion models in network optimization," *IEEE Commun. Surv. & Tut.*, vol. 26, no. 4, pp. 2611–2646, Fourthquarter 2024.
- [89] T. Wu *et al.*, "CDDM: Channel denoising diffusion models for wireless semantic communications," *IEEE Trans. Wireless Commun.*, vol. 23, no. 9, pp. 11 168–11 183, Sept. 2024.
- [90] H. Pan *et al.*, "CLEAR: Channel learning and enhanced adaptive reconstruction for semantic communication in complex time-varying environments," *arXiv preprint, arXiv:2412.08978*, 2024.
- [91] Z. Jin *et al.*, "GDM4MMIMO: Generative diffusion models for massive MIMO communications," *arXiv preprint, arXiv:2412.18281*, 2024.
- [92] B. Fesl, M. Baur, F. Strasser, M. Joham, and W. Utschick, "Diffusion-based generative prior for low-complexity MIMO channel estimation," *IEEE Wireless Commun. Lett.*, vol. 13, no. 12, pp. 3493–3497, Dec. 2024.
- [93] W. Lin, Y. Yan, L. Li, Z. Han, and T. Matsumoto, "SemantIC: Semantic interference cancellation towards 6G wireless communications," *IEEE Commun. Lett.*, vol. 28, no. 8, pp. 59–63, Aug. 2024.
- [94] D. Gündüz, M. A. Wigger, T.-Y. Tung, P. Zhang, and Y. Xiao, "Joint source-channel coding: Fundamentals and recent progress in practical designs," *arXiv preprint, arXiv:2409.17557*, 2024.
- [95] E. Boursoulatz, D. B. Kurka, and D. Gündüz, "Deep joint source-channel coding for wireless image transmission," *IEEE Trans. Cog. Commun. & Netw.*, vol. 5, no. 3, pp. 567–579, Sept. 2019.
- [96] D. B. Kurka and D. Gündüz, "DeepJSCC-f: Deep joint source-channel coding of images with feedback," *IEEE J. Select. Areas Inf. Theory*, vol. 1, no. 1, pp. 178–193, May 2020.
- [97] K. K. Wong *et al.*, "Virtual FAS by learning-based imaginary antennas," *IEEE Wireless Commun. Lett.*, vol. 13, no. 6, pp. 1581–1585, Jun. 2024.
- [98] J. Ho, A. Jain, and P. Abbeel, "Denoising diffusion probabilistic models," in *Proc. Int. Conf. Neural Inf. Proc. Syst. (NIPS)*, no. 574, pp. 6840–6851, 6–12 Dec. 2020, Vancouver BC, Canada.
- [99] E. Nachmani, R. S. Roman, and L. Wolf, "Non Gaussian denoising diffusion models," *arXiv preprint, arXiv:2106.07582*, 2021.
- [100] H. Guo *et al.*, "Gaussian mixture solvers for diffusion models," in *Proc. Int. Conf. Neural Inf. Process. Syst.*, 10–16 Dec. 2023, New Orleans, USA.
- [101] H. Wang, D. Zhai, X. Zhou, J. Jiang, and X. Liu, "Mix-DDPM: Enhancing diffusion models through fitting mixture noise with global stochastic offset," *ACM Trans. Multimedia Comput. Commun. Appl.*, vol. 20, no. 9, Sept. 2024.
- [102] O. Ronneberger, P. Fischer, and T. Brox, "U-Net: Convolutional networks for biomedical image segmentation," in *Medical Image Comput. & Computer-Assisted Intervention – MICCAI 2015*, pp. 234–241, 5–9 Oct. 2015, Munich, Germany.
- [103] M. K. Simon and M.-S. Alouini, "Digital Communication over Fading Channels," 2nd Ed., Wiley-Interscience, 2005, Chapter 2, pp. 41–47.
- [104] A. Papoulis and S. U. Pillai, "Probability, Random Variables and Stochastic Processes," 4th Ed., McGraw-Hill, 2002, Appendix C.

- [105] F. R. Ghadi *et al.*, "Fluid antenna multiple access with simultaneous non-unique decoding in strong interference channel," *IEEE Trans. Wireless Commun.*, doi:10.1109/TWC.2025.3578280, 2025.



Noor Waqar received his B.E. degree in Electrical Engineering from National University of Sciences and Technology (NUST), Pakistan, in 2021, and is currently pursuing his Ph.D. degree from University College London. His research interests include 5G and 6G communication, information theory, fluid antenna systems (FAS), MIMO communications, non-orthogonal multiple access (NOMA), and network optimization using machine and reinforcement learning techniques for wireless communication.



Kai-Kit Wong (Fellow, IEEE) received the B.Eng., M.Phil., and Ph.D. degrees in electrical and electronic engineering from The Hong Kong University of Science and Technology, Hong Kong, in 1996, 1998, and 2001, respectively. After graduation, he took up academic and research positions at The University of Hong Kong; Lucent Technologies; Bell-Laboratories; Holmdel; the Smart Antennas Research Group, Stanford University; and the University of Hull, U.K. He is the Chair of wireless communications with the Department of Electronic

and Electrical Engineering, University College London, U.K. His current research centers around 6G and beyond mobile communications. He is a fellow of IEEE and IET. From 2020 to 2023, he was the Editor-in-Chief for IEEE Wireless Communications Letters.



Chan-Byoung Chae (Fellow, IEEE) received the Ph.D. degree in electrical and computer engineering from The University of Texas at Austin (UT), USA in 2008, where he was a member of wireless networking and communications group (WNCG). Prior to joining UT, he was a Research Engineer at the Telecommunications R&D Center, Samsung Electronics, Suwon, South Korea, from 2001 to 2005. He is currently an Underwood Distinguished Professor and Lee Youn Jae Fellow (Endowed Chair Professor) with the School of Integrated Technology, Yonsei University, South Korea. Before joining Yonsei, he was with Bell Labs, Alcatel-Lucent, Murray Hill, NJ, USA, from 2009 to 2011, as a Member of Technical Staff, and Harvard University, Cambridge, MA, USA, from 2008 to 2009, as a Post-Doc. Fellow and Lecturer.

Dr. Chae was a recipient/co-recipient of the Korean Ministry of ICT and Science Award in 2024, the Korean Ministry of Education Award in 2024, the KICS Haedong Scholar Award in 2023, the CES Innovation Award in 2023, the IEEE ICC Best Demo Award in 2022, the IEEE WCNC Best Demo Award in 2020, the Best Young Engineer Award from the National Academy of Engineering of Korea (NAEK) in 2019, the IEEE DySPAN Best Demo Award in 2018, the IEEE/KICS Journal of Communications and Networks Best Paper Award in 2018, the IEEE INFOCOM Best Demo Award in 2015, the IEIE/IEEE Joint Award for Young IT Engineer of the Year in 2014, the KICS Haedong Young Scholar Award in 2013, the IEEE Signal Processing Magazine Best Paper Award in 2013, the IEEE ComSoc AP Outstanding Young Researcher Award in 2012, and the IEEE VTS Dan. E. Noble Fellowship Award in 2008.

Dr. Chae has held several editorial positions, including Editor-in-Chief of the IEEE Transactions on Molecular, Biological, and Multi-Scale Communications, Senior Editor of the IEEE Wireless Communications Letters, and Editor of the IEEE Communications Magazine, IEEE Transactions on Wireless Communications, and IEEE Wireless Communications Letters. He was an IEEE ComSoc Distinguished Lecturer from 2020 to 2023 and is an IEEE VTS Distinguished Lecturer from 2024 to 2025. He is an elected member of the National Academy of Engineering of Korea.



Ross Murch (S'84-M'90-SM'98-F'09) received the bachelor's and Ph.D. degrees in Electrical and Electronic Engineering from the University of Canterbury, New Zealand. He is currently a Chair Professor in the Department of Electronic and Computer Engineering (ECE) at the Hong Kong University of Science and Technology (HKUST). His research focus is creating new RF wave technology for making a better world. His unique expertise lies in his combination of knowledge from both electromagnetics and wireless communication systems and he

has published widely in both areas. He became a Fellow of IEEE for his "Contributions to Multiple Antenna Systems for Wireless Communication" in 2010. He is also a Fellow of IET, HKAE, HKIE, NASI and in 2023 he won the IEEE Communications Society Wireless Technical Committee Recognition Award. He is the founding Editor-in-Chief of the IEEE Journal of Selected Topics in Electromagnetics, Antennas and Propagation which was launched in January 2025. He was also an Area Editor for IEEE Transactions on Wireless Communications and Editor for IEEE Journal of Selected Areas in Communication: Wireless series (the precursor to IEEE Transactions on Wireless Communications). He has also been a guest editor for IEEE Proceedings and IEEE Journal of Selected Areas in communication. Professor Murch also enjoys teaching and has won several teaching awards. He has served IEEE in several roles and has been Department Head of ECE@HKUST from 2009-2015.